



Deposited via The University of Sheffield.

White Rose Research Online URL for this paper:

<https://eprints.whiterose.ac.uk/id/eprint/217109/>

Version: Published Version

Article:

Miltner, K.M. (2024) "A.I. is holding a mirror to our society": Lensa and the discourse of visual generative AI. *Journal of Digital Social Research*, 6 (4). pp. 13-33. ISSN: 2003-1998

<https://doi.org/10.33621/jdsr.v6i440456>

Reuse

This article is distributed under the terms of the Creative Commons Attribution-ShareAlike (CC BY-SA) licence. This licence allows you to remix, tweak, and build upon the work even for commercial purposes, as long as you credit the authors and license your new creations under the identical terms. All new works based on this article must carry the same licence, so any derivatives will also allow commercial use. More information and the full terms of the licence here: <https://creativecommons.org/licenses/>

Takedown

If you consider content in White Rose Research Online to be in breach of UK law, please notify us by emailing eprints@whiterose.ac.uk including the URL of the record and the reason for the withdrawal request.

“A.I. is holding a mirror to our society”

Lensa and the discourse of visual generative AI

Kate M. Miltner

University of Sheffield, United Kingdom

✉ k.miltner@sheffield.ac.uk

Abstract

Since the release of the open-source AI model Stable Diffusion in August 2022, a panoply of apps that create AI-generated portraits, or avatars, have exploded in popularity. One of the most notable examples is AI-powered photo editing app Lensa, owned by Prisma Labs. Lensa launched in 2018 but went viral in late 2022 due to the launch of its Magic Avatar feature, which uses Stable Diffusion to create fantastical (and occasionally bizarre) portraits of users. This paper analyzed the global English-language press coverage of Lensa and found that it focused on the app’s predatory data practices, the biased content it produced, and the user behaviors associated with it. I argue that this coverage provides evidence of discursive closure around key issues associated with visual generative AI that supports the maintenance of the status quo. I suggest that the press coverage of Lensa, which both articulates key AI-related harms and frames those harms as intractable and insolvable, creates a discourse of inevitability that has implications for how these issues are understood by the public, and for the approaches that are taken to address them. In doing so, it not only offers distinct advantages to those who stand to benefit most from this discourse, but also forecloses more imaginative public discussions of what visual generative AI could– or should– be.

Keywords: generative AI; discursive closure; selfies; algorithmic bias

1. Introduction

“Beware a world where artists are replaced by robots. It’s starting now” intoned a Los Angeles Times headline (Crabapple, 2022). “The inherent misogyny of AI portraits – Amelia Earhart rendered naked on a bed”, proclaimed The Guardian (Demopolous, 2022). “Your selfies are helping AI learn. You did not consent to this”, cautions The Washington Post (Ovide, 2022). All three of these headlines offer different warnings about the dangers of AI, and all three headlines are about the same app: Lensa, a photo editing app that makes AI-powered selfies.

Since the release of the open-source AI model Stable Diffusion in August 2022, a panoply of apps that create AI-generated portraits, or avatars, have exploded in popularity. Lensa, one of the most notable examples of this app category, originally launched in 2018 but went viral in late 2022 due to the addition of its Magic Avatar feature. Lensa uses the Stable Diffusion AI model to create fantastical (and occasionally bizarre) portraits of users, which are created by uploading a series of photographs and paying

a small fee; the app then generates anywhere from 50 to 200 AI-generated images in a series of user-selected themes, such as “Rock Star”, “Fairy Princess”, “Superhero”, and “Astronaut” (see Figure 1). After the Magic Avatar launch in November 2022, Lensa was installed over 13.5 million times worldwide, with users spending over USD\$29 million on the app in 12 days (McCluskey, 2022). In a matter of weeks, Lensa was the number one app in the Apple App Store’s Photo & Video charts, beating out heavyweights like YouTube and Instagram (Silberling, 2022). Lensa’s meteoric rise in popularity can be largely attributed to the fact that many users—including some public figures—posted their Magic Avatars on social media, sparking curiosity and a flood of global press coverage that both hyped and critiqued the phenomenon.



Figure 1. Examples of the author’s Magic Avatars

Despite the ostensible novelty of Magic Avatars, the press coverage about Lensa recycles common concerns about emerging technologies, particularly around algorithmic bias, unethical data use, and the narcissism of digital self-representation. While some of Lensa’s press coverage simply described the app and how to use it, a significant number of articles were highly critical, suggesting that the app engaged in predatory data practices, produced biased content, and encouraged user behaviors that were frivolous at best and harmful at worst. However, while the discourse about Lensa’s Magic Avatars articulates key AI-related harms, it also treats those harms as unfortunate but intractable problems that are already too difficult to solve. This suggests that some discursive closure (Deetz, 1992; Leonardi & Jackson, 2003; Markham, 2021) has already developed around both the risks and the potential of visual generative AI, which has implications for how these technologies are understood by the public. By relying on well-worn understandings of the key concerns raised by AI and other algorithmic technologies, the press coverage of Lensa creates a “discourse of inevitability” (Leonardi & Jackson, 2003; Markham, 2021) that makes these harms seem inescapable and unaddressable.

While Lensa’s moment was seen to be “over” (Lovejoy, 2023) in early 2023—indeed, the AI avatar app trend was said to have had “fizzled” (Perez, 2023) after a few months—these discourses remain as regular features in the press as competitor apps continue to launch and their features are hyped and occasionally dissected (e.g., Brown, 2023; Khan, 2023). As such, analyzing the discourse around Lensa offers an opportunity to reflect on the broader processes of AI publicization and their attendant implications for “exactly what AI is supposed to represent in the world” (Broussard et al., 2019). By promoting a discourse of inevitability around the harms of visual generative AI, the press coverage of Lensa offers “a particular view of reality that is maintained at the expense of equally plausible ones” (Deetz, 1992, p. 188); this ends up reinforcing narrow understandings of the issues at hand. In doing so, it not only contributes to the calcification of the broader discourse around AI and its “myths” (Natale and Ballatore, 2020) but also forecloses more imaginative public discussions of what visual generative AI could— or should— be.

2. The hype and harms of visual generative AI

Generative AI is a colloquial term for deep-learning computer models that can take a collection of data as an input (e.g., an artist's full body of work) and analyze that data to generate statistically probable outputs when provided with a prompt; these models "encode a simplified representation of their training data and draw from it to create a new work that's similar, but not identical, to the original data" (Martineau, 2021). In other words, generative AI models "find patterns in the data they are trained on and then create new work by attempting to mimic those patterns" (Miltner and Highfield, 2024).

There are a variety of generative AI models that can generate and/or modify a wide variety of media, including text, images, speech, video, and audio. Large language models (LLMs) – such as ChatGPT—produce text, and diffusion models are most often used in audiovisual modification and generation. Some of the most well known image-generating diffusion models include Midjourney, Firefly (Adobe), DALL-E (OpenAI), Llama (Meta), and Stable Diffusion (Stability AI). While many dominant models are proprietary, Stable Diffusion is an open-source model that was released to the public in August 2022 (Stability AI, 2022) and is used by a variety of generative AI tools, including Lensa (Prisma Labs, n.d.).

There has been considerable hype around visual generative AI tools, with claims that they are augmenting (Eapen et al., 2023), reimagining (Adobe Communications Team, 2023), and transforming (Roose, 2022) creativity and creative work. Indeed, the application of visual generative AI tools across a variety of fields including the visual arts, advertising, and film & TV has caused both celebration (e.g., Roose, 2022) and consternation (Equity, n.d.; Scherer, 2024). While some uses of visual generative AI tools have been decied as "inappropriate, unprofessional and disrespectful to audiences" (Sun and Harmon, 2024), others have been heralded as uses of "techno-vernacular creativity" that facilitates creative engagement within historically marginalized communities (Gaskins, 2022).

However, like many other AI technologies (e.g., Gillespie, 2024), visual generative AI models have been found to reproduce a series of harmful biases related to any number of identity characteristics, including gender (Sandoval-Martin and Martinez-Sanzo, 2024), race (Offert and Phan, 2022), mental health disorders (King, 2022), and religion (Alfano et. al, 2024). Quite often, these biases and stereotypes are imbricated with each other: Bianchi et al. (2023) found that Stable Diffusion generated "dangerous racial, ethnic, gendered, class, and intersectional stereotypes" that illustrated the "vast potential" of these models for "propagating harm along many axes of demographic identity" (p. 1494).

While it can be challenging to identify the root cause of these biases and stereotypes within a diffusion model (Luccioni et al., 2023), a model's training data has a significant impact on what it can do and how successfully it can do it (Buschek and Thorp, n.d., see also Denton et al., 2021). Midjourney, Stable Diffusion, and "perhaps hundreds" of other commercial models have been trained on LAION-5B (Buschek and Thorp, n.d.), an "uncurated" open-source dataset of over 5 billion image-text pairs that contains "strongly discomfoting and disturbing content for a human viewer" (Beaumont, 2022) including child sex abuse material (Thiel, 2023). An earlier version of this dataset (LAION-400) was found to have "troublesome and explicit images and text pairs of rape, pornography, malign stereotypes, racist and ethnic slurs, and other extremely problematic content" (Birhane, Prabhu, and Kahembwe, 2021). While LAION offers a disclaimer that LAION-5B should not be used in "ready to go industrial products", this warning has been largely ignored (Buschek and Thorp, n.d.) with clear consequences, as the Lensa case study illustrates.

3. Controversies & closures in AI media coverage

Scholarship on the influence of the news media has long established that media coverage about technology heavily influences how the public understands emerging technologies, as well as their risks and benefits; this is particularly true when audience members may not be familiar with a particular issue (Scheufele and Lewenstein, 2005). Chuan et al. (2019) argue that research evaluating how the media

discusses AI is essential to understand how public opinion about it is shaped; furthermore, as Brause et al. (2023) have noted, the media plays an important role in the dissemination of information about AI, shaping public perceptions and influencing public debates where “certain understandings of AI are advanced, and different stakeholders critically engage with the technology and its role in society, contributing to the communicative construction of the technology” (p.1). Ouchchy et al (2020) explain that it is particularly important to understand how the media portrays issues of AI and ethics, as they could have significant implications for AI development and regulation:

Because the members of the general public, as both consumers in the market economy and constituents of a liberal democracy, are key stakeholders for technology adoption—and, to a certain extent, for public policy and regulatory oversight—public opinion could affect what kind of AI is developed in the future and how AI is regulated by the government. (p. 927)

The representation of AI in the global media is a developing area of research, with approximately 30 empirical studies published since 2017 (Brause et al., 2023). These studies suggest that AI is usually discussed in the context of business or economics, and that the valence of this coverage is overwhelmingly positive (Brause et al., 2023; Chuan et al., 2019; Fast and Horvitz, 2017). In this sense, the media discourse surrounding Lensa is a bit of an anomaly, as it focuses more on social issues and includes a significant amount of critical coverage. However, as Fast and Horvitz (2017) and Chuan et al (2019) have pointed out, concerns about AI have also increased in news coverage in recent years, particularly in relation to AI ethics, the negative impact of AI on work, privacy, and loss of control over AI.

Dandurand et al (2023) note that the media has “become a powerful site of discourse formation in which certain voices and tropes about AI are authoritatively put forth while others are not” (p.2). Through this process, the media “participates in making, or not making more exactly, AI controversial” (p.3). Marres (2015) notes that in Science and Technology Studies (STS), controversy studies illustrates the cozy relationship between the formulation of knowledge claims and the organization of political interests (p. 656). The news coverage of Lensa certainly appears controversial in its discussion of data exploitation, misogyny, and racism in particular; these hot-button issues are often framed in a polemical way that seems to foment public discussion about the treatment of marginalized groups within algorithmic systems. However, what is notable about the Lensa coverage is that instead of providing novel pathways for public debate about these concerns and how to address them, generative AI harms are discussed in a way that is fossilized, in that it is frozen in place and stuck in the past (Dihal, 2024). Dandurand et al. (2023) ask, “what to do when AI’s coverage involves so much consensus about its benefits?”; conversely, the question that the Lensa coverage raises is: what to do when AI’s coverage involves so much consensus about its harms?

It is here where the communication studies concept of discursive closure can provide some assistance. Derived from Habermas’s concept of systematically distorted communication, Stanley Deetz (1992) theorizes discursive closure as existing “whenever potential conflict is suppressed” (p. 187). Discursive closure can take place through a variety of communicative processes and practices that “shut down or close off options for thinking otherwise” (Markham, 2021, p. 392). Such processes include disqualification (the exclusion of individuals from a discussion or conversation), naturalization (the reification of social dynamics and structures), neutralization (the positioning of value-laden activities as if they were neutral), and topical avoidance (Deetz, 1992; Leonardi and Jackson, 2003).

Markham (2021) suggests that discursive closure allows us to see how “certain patterns of thought, talk, actions, or interactions tend to function like negative feedback loops in social ecologies, discouraging evolution and change” (p. 392). Deetz (1992) argues that one way that discursive closure is achieved is “through the privileging of certain discourses and the marginalization of others” (p. 187), and that such discursive practices can either contribute to “further exploration of the subject matter” or “divert, distort, or block the open development of understanding” (p. 188). This curtailment of understanding and open discussion often has significant implications and consequences, usually in maintaining the status

quo: Deetz argues that when discussion of a particular issue is curtailed, “a particular view of reality is maintained at the expense of equally plausible ones, usually to someone’s advantage.” (p. 188)

Leonardi and Jackson (2003) have illustrated how discursive closure can deliver economic, political, and social advantages to certain ideologies over others by eliminating alternative conceptualizations of how things could have taken place or been done (p. 615). In their study of high-tech mergers in the late 1990s, they show how leaders of high-tech organizations invoked the “inevitability” of technology to justify managerial decisions to the public. Markham (2021) also found that discourses of inevitability were present in how members of the public engaged conceptually with digital platforms and algorithmic technologies, where alternatives to the status quo seemed “unimaginable” and any radical alternatives dismissed as improbable. She argues that discourses of inevitability teach us “that our present and possible futures are being determined by technology” and that even if the current trajectory suggests that “the future world is likely to be dystopian”, we are devoid of choice, because we are faced with an “either/or proposition” where we are either connected to technology or we do not exist (p. 397). As will be illustrated below, the theme of inevitability was both implicit and explicit in the discourse surrounding Lensa, foreclosing any significant debate or discussion of what AI should be and how it should operate.

4. Research design & methodological considerations

The study corpus consisted of 135 unique articles from the global English-language syndicated and tech press (see Figure 2). Articles were collected through two complementary methods. First, I searched Nexis for the term “Lensa AI” in English-language coverage from November 2022 to July 2023 across the system-defined regions of Africa, Asia, Europe, Latin America, Middle East, North America, and Oceania. This returned just under 400 results from 12 countries: Australia, Canada, Egypt, India, Mexico, Nigeria, Pakistan, Singapore, Switzerland, United Arab Emirates, United Kingdom, and United States. These results were refined by excluding content that was topically irrelevant. I then eliminated any duplicated or syndicated articles—of which there were many—to ensure that each article was only represented once in the corpus. This deduplication eliminated some countries from the analysis (i.e., Nigeria). After eliminating any irrelevant, duplicated, or syndicated articles, I was left with a total of 83 articles from Nexis. I then took a “snowball” approach (see Braithwaite, 2016) to the content linked or referred to in the Nexis results, adding any new articles that were specifically about Lensa. I then read these articles and followed any links in those pieces until I reached saturation and failed to come across any new material. I added a total of 52 articles through this process.

I analyzed the corpus using reflexive thematic analysis (Braun & Clarke, 2019; Byrne, 2020). I created a series of initial codes that reflected content that appeared repeatedly in the corpus; some examples of these codes include “skin tone”, “cleavage”, “celebrity”, “selfies”, “nudity”, “artist concerns”, “plastic surgery” and so on. These codes were then combined into broader themes, such as “sexualisation”, “privacy”, “narcissism”, and “training data”; these themes were then further condensed into the three overarching themes discussed below. I also reflected on the types/categories of coverage that were present in the sample: as will be discussed in the analysis, there were two types of coverage: coverage that focused on explaining what Lensa was (explainer coverage), and coverage that focused on critiquing Lensa (scandal coverage). After the topical themes and coverage categories were established, I returned to the corpus and reviewed each article, tallying which themes appeared within it and noting what coverage category it belonged to. This allowed me to note the prevalence of each theme and type of coverage within the corpus overall as well as any topical trends across countries/regions. It is worth noting that while articles were labelled as being predominantly one coverage type (explainer or scandal), some articles reflected multiple themes.

Table 1. Article totals

Total number of articles by country, as well as representative publications from each nation

	<i>Totals</i>	<i>Percentages</i>	<i>Sample Outlets</i>
<i>USA</i>	<i>46</i>	<i>34%</i>	<i>Techcrunch, Wired, MIT Technology Review, The New York Times</i>
<i>Mexico</i>	<i>42</i>	<i>31%</i>	<i>CE Noticias Financieras</i>
<i>UK</i>	<i>19</i>	<i>14%</i>	<i>The Independent, The Sun, The Guardian, Yorkshire Evening Post</i>
<i>India</i>	<i>17</i>	<i>13%</i>	<i>The Times of India, The Telegraph, Hindustan Times, Indian Express</i>
<i>Australia</i>	<i>3</i>	<i>2%</i>	<i>The Age, The Guardian</i>
<i>Switzerland</i>	<i>2</i>	<i>1%</i>	<i>Handelszeitung</i>
<i>Canada</i>	<i>2</i>	<i>1%</i>	<i>Toronto Star</i>
<i>UAE</i>	<i>1</i>	<i>1%</i>	<i>Wknd Magazine</i>
<i>Singapore</i>	<i>1</i>	<i>1%</i>	<i>The Straits Times</i>
<i>Egypt</i>	<i>1</i>	<i>1%</i>	<i>The Egypt Independent</i>
<i>Pakistan</i>	<i>1</i>	<i>1%</i>	<i>Pakistan and Gulf Economist</i>
<i>Total</i>	<i>135</i>	<i>100%</i>	

Nexis indicated that there were a few hundred results for “Lensa AI” in languages including Turkish, Spanish, French, German, Dutch, Arabic, and Portuguese. A cursory review of the French and Spanish language results suggests that some of the themes that were present in the English coverage are also present in other languages. At least part of the reason for this is because some English language news sources are translated into other languages and syndicated: for example, I found a Spanish-language article that turned out to be a translated version of a WIRED article that was already in the corpus. Of course, it is possible that there are themes present in other languages that aren’t reflected in this analysis, but it was notable that much of the coverage in this corpus was syndicated, particularly from US-based sources. This was also the case in other English-speaking countries: in Australia, for example, I only found three unique articles, two of which were from The Guardian; even when searching directly on key news sites like the Sydney Morning Herald to ensure I hadn’t missed anything significant, I only found syndicated materials from e.g., The Washington Post and the Wall Street Journal that were already included in the corpus. Furthermore, it is worth noting that there were more articles (and more countries) represented in the corpus before I traced syndicated articles to their sources and eliminated the duplicates. This suggests that the discourse about Lensa likely has a wider global audience than the countries reflected in this analysis, although the syndicated content clearly reflects the material included in this analysis. The issue of syndication and other factors related to the global news industry will be addressed more thoroughly in the Conclusion, but it is also a methodological consideration for this study.

Overall, my goal with this analysis was to offer an “illustrative rather than exhaustive account” (Braithwaite, 2016), of the discourse about Lensa, demonstrating its most salient features (p. 3). As the next section will illustrate, the discourse surrounding Lensa had distinct contours that appeared with

notable consistency across the countries represented in the corpus, albeit in different proportions. This clarity and consistency suggest that this analysis offers a clear illustration of the key issues associated with visual generative AI— and how they are discussed in the press— at this particular moment.

5. “Sex, art theft, and privacy”: Lensa in the global media

Although there were 11 countries represented in the corpus, there were four countries that comprised the majority of English-language Lensa coverage: the United States, Mexico, the United Kingdom, and India. The United States and Mexico dominated the coverage, with 34% and 31% of the corpus, respectively. News items from the United Kingdom comprised 14% of the sample, and Indian stories were 13%. Collectively, these four countries made up 92% of the corpus. The remaining 7 countries (Australia, Canada, Egypt, Pakistan, Singapore, Switzerland, and the UAE) collectively composed 8% of the corpus. While the dominance of the US, Mexico, the UK, and India within the sample likely has many underlying factors, one aspect could be related to the popularity of the app in each country: Lensa was the most downloaded app in the US in December 2022 (Silberling, 2022), and it was 6th on the charts in India (Rekhi, 2022).

The Lensa discourse had an intense burst of activity in the first two weeks of coverage, followed by a long tail: Lensa launched its Magic Avatar feature in late November 2022, the first article in the corpus was published in Mexico on November 29, 2022, and the last article in the corpus appeared in India on July 27, 2023.¹ However, 81% of the coverage in the corpus (110 articles) appeared within a month from when the first article was published; given that the corpus does not reflect any duplicate or syndicated content (which would increase those numbers considerably), that is a significant amount of coverage in a condensed period of time. As the chart below indicates, there were two spikes in coverage during this time: the first was on December 8, shortly after articles in TechCrunch and WIRED revealed that Lensa was generating explicit content, and the second spike was on December 13, the day after Melissa Heikkila (2022b) from MIT Technology Review published an article discussing the sexualized and racialized nature of her Lensa avatars.

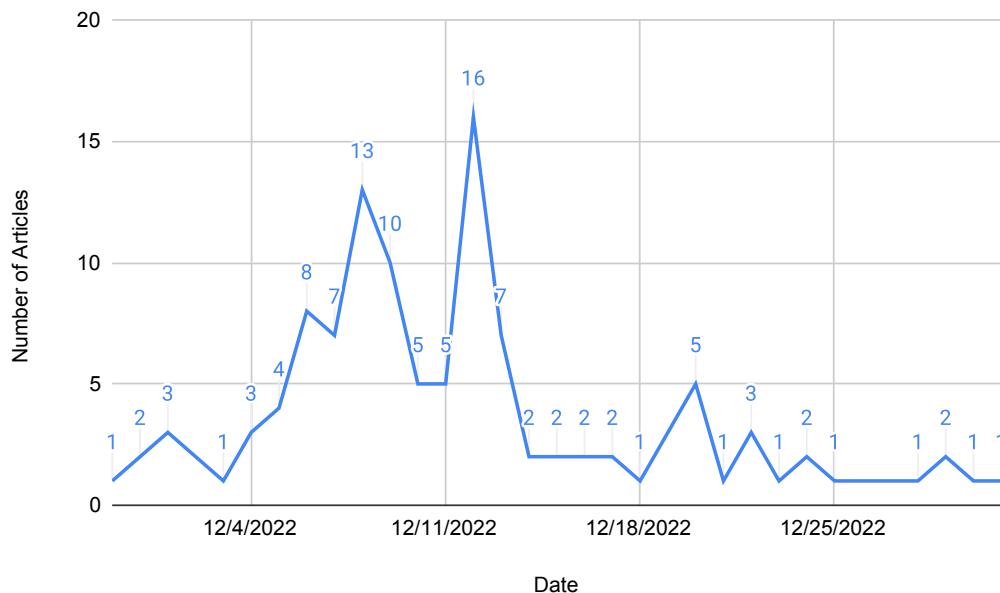


Figure 2. Article frequency by publication date

¹ It is possible that there are more articles that have appeared after July 2023, but that is the point at which data collection concluded

As mentioned in the Methodology section, there were two distinct types, or genres, of content in the Lensa discourse: scandal coverage and explainer coverage. Scandal coverage, which comprised 42% of the corpus, focused on perceived wrongdoing or norm violations related to the app and its use. Scandal coverage began to appear in the corpus on December 5. While scandal coverage was found in all countries aside from Egypt, Pakistan, UAE, and Singapore,² the topics that dominated the scandal-related articles came from a handful of articles (e.g., Hatmaker, 2022; Heikkila, 2022b; Kamps, 2022; Snow, 2022) published in US-based, agenda-setting tech publications like TechCrunch, MIT Technology Review, and WIRED. These articles were published within the first two weeks of the news cycle and were very influential; not only did they shape the discourse on a global level, but they were frequently linked to in other publications.

The other type of coverage was explainer coverage, which constituted 58% of the corpus and focused on describing what Lensa is and how it worked. Explainer coverage was present in the corpus from November 29th; indeed, all of the articles from the first six days of the news cycle fell into this category. Explainer coverage typically came from mainstream news publications. It contextualized Lensa as a viral phenomenon, talking about how “everyone” (or in some cases, specific celebrities) were using it. It sometimes included explanations of other generative AI apps, whether they were image-based (e.g., Remini, Midjourney, Firefly, DALL-E) or text-based (e.g., Chat GPT, Otter, Bard). This coverage typically explained the Magic Avatar feature, how to download and use the app, how much it cost, and if the results were any good. Some explainer coverage did this in a completely uncritical way, praising Lensa’s innovative nature and talking about how “cool” the generated images were. Most explainer coverage, however, tended to summarize or mention the key issues referenced within the scandal coverage as part of its description of the app as a hot new trend.

As noted previously, the discourse is remarkably consistent across contexts in terms of its content; in addition to the two categories of coverage, there were three overarching themes across the sample. The first theme focused on Predatory Data Practices. Articles related to this theme discussed unethical, exploitative practices relating to the non-consensual collection and use of training data, particularly from artists and app users. The second theme was Biased Content. Articles relating to this theme referred to the gender and racial biases reflected in some of the images generated by Lensa, including the sexualization of women, the lightening of skin tones, and the generation of explicit content. The third theme related to User Behavior. The content in this theme described how people used the app in specific ways, including celebrities and political figures. Below is a table that illustrates the distribution of each type of coverage and theme within the sample (Table 2).

² The lack of scandal coverage in these countries is likely attributable to the fact that there was only one article in the corpus for each of them.

Table 2. Article Distribution*Distribution of coverage type and themes across countries within the sample*

	USA	%	MEX	%	UK	%	IND	%	AUS	SWZ	CAN	EGY	PAK	UAE	SIN	Total	Total %
<i>Explainer</i>	19	41%	32	76%	7	37%	14	82%	1	2	0	1	1	1	1	79	59%
<i>Scandal</i>	27	59%	10	24%	12	63%	3	18%	2	0	2	0	0	0	0	56	41%
	46		42		19		17		3	2	2	1	1	1	1	135	
<i>Predatory Data Practices</i>	35	76%	15	36%	12	63%	9	53%	3	0	1	0	0	0	0	75	56%
<i>Biased Content</i>	30	65%	8	19%	12	63%	9	53%	1	2	1	0	0	0	0	63	47%
<i>User Behavior</i>	18	39%	18	43%	9	47%	4	24%	1	0	0	0	0	0	0	50	37%

5.1 Predatory data practices

The Predatory Data Practices theme was the most prevalent theme in the corpus, reflected in 56% of the articles and within six countries (USA, Mexico, UK, India, Australia, and Canada). It was most commonly discussed in the US and UK articles (76% and 63%, respectively), with a little more than half of the Indian articles (53%) and one-third of the Mexican articles (36%) referencing data-related issues.

This theme focused predominantly on two distinct, but related topics: the use of art as training data in the Stable Diffusion model and Lensa’s privacy policy, which stated at the time that users’ photos could be used to train the app. Coverage reflecting this theme largely focused on the unethical or non-consensual use of data in AI technologies and the implications of those practices for both specific communities and users in general. While these practices were heavily critiqued for the distress and harm that they caused, they were also largely framed as unavoidable. This aligns with Markham’s (2021) observation that there is an “overall expectation” that technology companies collect and sell (personal) data and that an alternative to this state of affairs seems unimaginable (pp. 390-391).

5.1.1 “Robot replacements”: AI and/as art

The discussion of art and artists within the Lensa discourse originated in the US, but was also the focus of articles in Canada, Australia, and the UK. It was not the first time a debate about art, the ethics of training data, and the devaluation of creative labor had appeared in the media; rather, the launch of Lensa “reignited discussion” (Sung, 2022) of these issues.

Within the corpus, there was a common framing that Lensa was “stealing” from artists. This phrase first appeared in the corpus in a TechCrunch (US) article by Taylor Hatmaker (2022) on December 5, noting that “for every 10 Lensa avatars there’s one Cassandra in the comments scolding everyone for paying for an app that steals from artists”. In an NBC News (US) article on December 7, (Sung, 2022), artist Karla Ortiz commented, “Companies like Lensa say they’re “bringing art to the masses, but really what they’re bringing is forgery, art theft [and] copying to the masses.” Within a week of Hatmaker’s original article, this topic appeared in articles across the globe.

The “theft” that Lensa was seen to perpetuate—described in one article as “arguably the biggest art heist in history” (Tran, 2022)—had three intertwined elements: the unethical use of art as training data, a lack of fair compensation to artists, and the replacement of artists by AI. This was discussed as a multi-layered insult: artists’ work was used without consent— and without remuneration— to train models which were duplicating their styles³ and stealing jobs from them.⁴ These concerns were summarized by artist Molly Crabapple (2022) an op-ed in the Los Angeles Times (US), who stated,

These data sets were not ethically obtained. LAION sucked up 5.8 billion images from around the internet, from art sites such as DeviantArt, and even from private medical records...They took it all without the creator’s knowledge, compensation, or consent. Once LAION had scraped up all this work, it handed it over to for-profit companies — such as Stability AI, the creator of the Stable Diffusion model — which then trained their AIs on artists’ pirated work... AIs can spit out work in the style of any artist they were trained on — eliminating the need for anyone to hire that artist again.

Similar claims were repeated throughout the corpus. In *The Independent* (UK), Josh Marcus (2022) noted that the AI models used by apps like Lensa “harvest the stylistic DNA of individual artists, then allow strangers to borrow elements from their work without offering any credit” in an “incredibly direct” manner. This was echoed by Australian artist Kim Leutwyler, who explained to *The Guardian* (AUS) that “some artists are having their exact style replicated exactly in brush strokes, colour, composition – techniques that take years and years to refine” (Kelly, 2022). While some articles also noted that AI models are seen by some artists as “a help, not a threat” (Marcus, 2022) and that something that “augments” the capabilities of artists (Flux, 2022), the majority of articles framed Lensa and its ilk primarily as “robot replacements” that “vampirized the work of artists” (Crabapple, 2022). Furthermore, this state of affairs was largely framed as a *fait accompli*, with artists being described as “exhausted and powerless” (Tran, 2022) with no recourse, because it is “unclear” how artists can protect their intellectual property “from being sucked into AI models” even if they try (Marcus, 2022). This sense of inescapable data collection was also present in the coverage that discussed Lensa’s privacy policy.

5.1.2 “You did not consent to this”: Lensa’s privacy policy

Concerns about Lensa’s privacy policy appeared around the same time as concerns about Lensa as an art “thief”: on December 5, articles in *Art News* (US) and *Refinery29* (UK) noted that Lensa’s terms and conditions allow it to “use the manipulated photos you upload for any way it sees fit” (Roberts, 2022). In an article titled “Careful — Lensa Is Using Your Photos to Train Their AI”, Shanti Escalante-DeMattei (2022) noted that “Lensa’s privacy policy and terms of use stipulate that the images users submit to generate their selfies, or rather the ‘Face Data,’ can be used by Prisma AI, the company behind Lensa, to further train the AI’s neural network.” The privacy policy was brought up again on December 7 at the end of an explainer article in *The New York Times* (Kircher and Holtermann, 2022), where it was explained that users’ “face data” was deleted within 24 hours, but that any uploaded photos or videos could be used to train Lensa’s “algorithm”. The issue of Lensa’s privacy policy appeared regularly after that, particularly in explainer coverage, where it was framed as an issue of “safety”, “risk” or “danger”, especially in India, Mexico, and the UK.

Similar to the art-related coverage, articles discussing Lensa’s privacy policy often made connections to the broader implications of user data being used to train AI systems. In the *New York Times* explainer article referenced above, Jen King, a privacy and data policy fellow at Stanford, suggested that it was user data that was Lensa’s profit motive, stating, “I doubt that the whole business model is, ‘Give us \$10 or \$15 and we’ll send you back an A.I. glam shot’” (Kircher and Holtermann, 2022). In a December 10 *Washington Post* article, Shira Ovide (2022) argued that Lensa and other “cool AI technologies” are “built on all the information we’ve put out into the world”, which is having significant unintended consequences: “that drunk selfie you posted on Instagram” is now “training fuel for an artificial

³ Fantasy landscape artist Greg Rutkowski has become a poster child for this issue, as his name became one of the most used prompts when Stable Diffusion was released in August 2022 (See Heikkilä 2022a).

⁴ For an example of this, see Maimann, 2022.

intelligence system that helps put an innocent person in jail.” While these outcomes are framed as unfortunate, they are also framed as unavoidable: as Ovide put it, “once you put digital bits of yourself or your loved ones online, you lose control of what happens next.”

Discussions of Lensa’s privacy policy also often harkened back to previous privacy scandals associated with apps like Faceapp (mentioned in US and Indian coverage), Snapchat filters, and Reface (both mentioned in UK coverage). This is not entirely surprising, given that before its Magic Avatar feature was launched, Lensa was exclusively an image-editing app. Image-editing apps were also invoked in discussions relating to specific types of aesthetics and beauty standards, such as lightened skin tones and normative body sizes. These kinds of biases were a significant topic of discussion on their own, particularly in the scandal coverage.

5.2 Biased Content

Biased Content was the second most prevalent theme in the corpus, reflected in 47% of articles. However, it was arguably the most impactful: the two spikes in coverage on December 8 and December 13 were related to articles that focused on this theme and had a big influence on its frequency in the coverage. The presence of the Biased Content theme was relatively even across US (65%), UK (63%) and Indian (53%) coverage. However, only 19% of the Mexican coverage referenced it, which featured predominantly in scandal coverage. The Biased Content theme also appeared in coverage from Australia and Canada, and the two articles from Switzerland focused exclusively on it.

The content in this theme largely focused on two topics: Lensa’s generation of explicit images that sexualized women and girls, and the fact that it created images that were racialized in specific ways. These topics were often intertwined, with some of the most influential articles engaging with both topics in a way that reflects the intersectionality of these concerns. The discussion of how women— and women from minoritised racial groups in particular— are represented by Lensa reflects a broader discourse about the biases of algorithmic technologies (e.g., Noble, 2018). However, unlike this discourse, which understands algorithmic bias as a multidimensional problem, the Lensa coverage laid the blame for Lensa’s production of explicit and biased imagery almost exclusively at the feet of the Stable Diffusion model and the “unfiltered internet content” (Prisma Labs, n.d.) that it was trained on. While the sexualized and racialized images Lensa produced were heavily critiqued by the press, they were also portrayed as an unfortunate— but perhaps unavoidable— reflection of the “biases humans incorporate into the images they produce” (Prisma Labs, n.d.).

5.2.1 “Why do all my AI avatars have huge boobs?”⁵: Explicit and sexualized content

The topic of explicit and/or sexualized content first appeared in the corpus on December 5 in a Refinery 29 (UK) article that noted that Twitter users had been complaining about the fact that Lensa had “slimmed down larger-bodied people” and perpetuated “the male gaze” through “hyperfeminised ideations of female-identifying individuals” (Roberts, 2022). The complaints about Lensa sexualizing women continued on December 6 in an op-ed in Business Insider (Hill, 2022) that explained how a Lensa user with a history of disordered eating was triggered by the “heroin chic” images that the app generated, making her appear “model-thin, sexualized, and miserable.”

The issue of sexualization was further amplified by articles in TechCrunch (US), WIRED (US), and The Guardian (UK) that discussed how Lensa produced explicit content. The TechCrunch article, titled “It’s way too easy to trick Lensa AI into making NSFW images”⁶ also came out on December 6. Author Haje Jan Kamps (2022) noted that, even though Stable Diffusion has an explicit content filter in place, images that have been modified by Photoshop somehow deactivate that filter and “gladly churn out a

⁵ Mercado, M. (2022, December 12). Why Do All My AI Avatars Have Huge Boobs? *The Cut*. <https://www.thecut.com/2022/12/ai-avatars-lensa-beauty-boobs.html>

⁶ NSFW is an acronym for “not safe for work”, aka explicit or pornographic content

number of problematic images”. He then highlighted the implications of being able to “create near-photorealistic AI-generated art images by the hundreds without any tools other than a smartphone, an app and a few dollars”, particularly in relation to the generation of AI-generated, or “deepfake”⁷ pornography, both for celebrities and normal people alike.

The following day, WIRED published an article titled “‘Magic Avatar’ App Lensa Generated Nudes From My Childhood Photos” (Snow, 2022). Researcher Olivia Snow (2022) detailed how, when she tested Lensa’s “no kids, adults only” policy by submitting a combination of “mousy” adolescent and “awkward” childhood photos, Lensa produced “fully nude photos of an adolescent and sometimes childlike face but a distinctly adult body.” Snow explained that she had been inspired to do this because many users (but “primarily women”) had pointed out that the app “ascribes cartoonishly sexualized features, like sultry poses and gigantic breasts, to their images” and generated “fully nude results despite uploading only headshots”, which in turn made people feel “very violated”. The Guardian (UK) ran a similar experiment, feeding Lensa images of “feminist icons” Amelia Earhart, Betty Friedan, and Shirley Chisholm, and receiving results comparable to Snow’s:

The author of *The Feminine Mystique* became a nymph-like, full-chested young woman clad in piles of curls and a slip dress. Chisholm, the first Black woman elected to US Congress, had a wasp waist. And the aviation pioneer was rendered naked, leaning on to what appeared to be a bed (Demopolous, 2022).

Lensa’s production of erotic imagery was primarily attributed to the LAION training dataset used by Stable Diffusion and other AI models. Articles in the Wall Street Journal (O’Brien, 2022) and MIT Technology Review (Heikkilä, 2022b) explained that because the internet is rife with images of naked or scantily clad women, the dataset that Stable Diffusion was trained on is skewed toward sexualized images, whether women want to be depicted in that manner or not. A spokesperson from Lensa’s parent company explained to the MIT Technology Review (Heikkilä, 2022b) that “the man-made, unfiltered online data introduced the model to the existing biases of humankind” and as such, they couldn’t possibly “consciously apply any representation biases or intentionally integrate conventional beauty elements” into the app and its output. In other words, the ultimate blame for Lensa’s production of sexualised imagery lies with humans and their biases, and not with the technology itself. This was an explanation that was also provided to explain Lensa’s production of images that reflected racial biases in addition to gender biases.

5.2.2 “Generic hot white girl”: Racialized images

The issue of racialization first appeared in the corpus in the same Art News article that highlighted the concerns with Lensa’s privacy policy (Escalante-DeMattei, 2022). It noted that writer Maya Kotomori had received a batch of avatars that made her look white, despite the fact that she is a “fair-skinned black person”. AI For The People founder Mutale Nkonde explained in The Washington Post (US) that the lack of representation of dark-skinned people in AI training images means that AI technologies like Lensa have trouble analyzing and reproducing images of them; as a result, Black women get images that look like a “generic hot white girl” (Hunter, 2023). This phenomenon was also described by Swiss writer Rebecca Stevens Alder (2022), who suggested that the “white privilege and racism” built into Lensa meant that it “did not know what to do when faced with a Black face”. Polygon’s Nicole Clark (2022), who is Taiwanese and white, similarly said that Lensa had “no idea what to do with my face” and created images that looked alternately East Asian or white, but like “complete strangers, save for one particular quirk or other of mine, like my jawline or eye shape”. Clark noted that whether she looked white or East Asian depended on the filter she chose, with the images for “Kawaii” looking more Asian and the images for “Light” and “Fantasy” looking white. She explained that while one image in the “Cosmic” set gave

⁷ Deepfakes are “hyper-realistic videos that apply artificial intelligence (AI) to depict someone say and do things that never happened” (Westerlund, 2019, p. 38); AI-generated images can also be considered deepfakes.

her “random cleavage”, she did not receive any nude photos; however, that was certainly not the case for other people with Asian ancestry.

The imbrication of sexualization and racialization in the results for people from minoritized racial groups was discussed at length in a viral article by Melissa Heikkila (2022b) of the MIT Technology Review that was published on December 12. Heikkila, who described herself as mixed race, stated that her Lensa avatars were “cartoonishly pornified” and noted that out of a total of 100 avatars, 14 were topless and a further 16 were in “skimpy” outfits and sexualized poses. She suggested that this extreme sexualization was attributable to her Asian heritage, which seemed to be the dominant feature that the AI model recognized.

I got images of generic Asian women clearly modeled on anime or video-game characters. Or most likely porn, considering the sizable chunk of my avatars that were nude or showed a lot of skin. A couple of my avatars appeared to be crying. My white female colleague got significantly fewer sexualized images, with only a couple of nudes and hints of cleavage. Another colleague with Chinese heritage got results similar to mine: reams and reams of pornified avatars. Lensa’s fetish for Asian women is so strong that I got female nudes and sexualized poses even when I directed the app to generate avatars of me as a male.

Snow (2022) also noted in her article that “a woman of Asian descent” explained to her that in photos where she didn’t look white, she was given “ahegao face”, a facial expression used in erotic manga to “highlight a hyperintense orgasm” (Santos, 2020).

Like the sexualization issue, the blame for Lensa’s racialization problems within this coverage is laid at the feet of existing structural inequalities: society is biased, consequently, so are the images that they create; if these are the images that AI models are trained upon, such content is inescapable and simply the cost of engagement. For users who receive inappropriate or biased content, what are the options? According to Glamour (UK)’s Anya Meyerowitz (2022), there is only one: don’t take part. “It’s clear that the app is violating our rights, perpetuating misogyny and is based on, at best, shaky ethics”, Meyerowitz concludes. “We’d encourage everyone, not just women, to stop using the app.”

5.3 User behavior

The idea that users should stop using Lensa is a common subtext in the User Behavior theme, although unlike the Predatory Data Practices and Biased Content themes, it is our own narcissistic tendencies– and not necessarily the AI models– that we should be eschewing.

The User Behavior theme encompassed content that discussed how specific people or groups were using Lensa. This included both average users and celebrities in both explainer and scandal content. User Behavior was the least prevalent theme in the American (39%), British (47%) and Indian (24%) coverage. It was, however, the most prevalent theme in Mexican coverage (43%), largely due to the volume of articles that mentioned how specific public figures shared their Magic Avatars on Instagram, including politicians (e.g., Mexican President Claudia Sheinbaum, Senator Ricardo Monreal) and celebrities (e.g., commentator and host Marisol Gonzalez, singer Yahaira Plasencia). Most of the User Behavior content in Mexico was explainer coverage, but there was also some Scandal content that talked about the influence of artificial intelligence on our self-perception. In an article titled “Looking like a cyborg to look pretty: how artificial intelligence is shaping the canon of beauty” (Bou, 2022), the author argues that the popularization of Lensa and other image-focused technologies promotes an unrealistic, homogenous, “robotic beauty” that reflects the subordination of our identities to the dictates of technological advances.

The relationship between identity and visual technologies like Lensa was a dominant thread in the User Behavior theme. An article in Time (US) suggested that there was a “vanity factor” in play to Lensa’s appeal, and a New York Times article titled “AI’s Best Trick Yet Is Showering Us With Attention” (Haigney, 2023) proposed that apps like Lensa are popular because they engage our narcissistic tendencies. “The fundamental appeal of apps like this is, of course, to our own self-involvement,” author Sophie Haigney argued, concluding that “the app tricks me into feeling seen, but really it is just me, trying once again to see myself.” An article in The Guardian (UK) (Demopolous, 2022) made a similar

argument, suggesting that the “allure” of Lensa was based in a self-oriented “curiosity”. Even though Lensa often produced images of “teenage-boy comic-book fantasy girls” that made her feel “icky”, illustrator Cat Willet explained that “even if you know it’s wrong, you still want to see what yours looks like”. While most articles framed this self-interest as a shallow, albeit relatively harmless, narcissism, other articles suggested that the “hot AI selfies” produced by Lensa might exacerbate body image concerns or disorders (Klee, 2022). Other articles suggested that engagement with Lensa might have even more extreme effects: an article from TMZ that was referenced in several articles in the corpus (“Celeb Surgeons Say”, 2022) reported that celebrity plastic surgeons were receiving visits from people requesting to look like their Lensa portraits.

There were some articles, however, that described Lensa being put to an empowering use: enabling transgender and non-binary people to experience “gender euphoria” through the production of images that align with their gender identity. In Refinery29 (UK), Millie Roberts (2022) discussed the “transformative power of avatars” for LGBTQ+ technology users to enhance “the lives of diverse and minority communities” by offering venues for identity exploration in online contexts. Adam Smith (2022) explained in Openly News (UK) that for users with gender dysphoria, their Lensa portraits could be “affirming” and Saskia Maxwell Keller (2022) of Out (US) described the “trans joy” inspired by Lensa as a “positive force”. However, Roberts, Smith, and Keller all noted that not all trans users had equally positive experiences, and that the controversies around Lensa’s training data made some trans users reluctant to engage with it.

Despite the relative novelty of Lensa’s Magic Avatar feature, the discourse that explains, analyzes, and critiques it does not offer a new perspective; instead, it treads well-worn paths that portray the negative aspects and consequences of visual generative AI as unavoidable or inevitable. The key themes of “sex, art theft and privacy” (Biron, 2023)—and, importantly, how they are discussed—in the Lensa discourse connect to existing framings and understandings of these issues. This has significant implications for how we understand the key problems inherent in AI technologies, as well as the solutions that we propose and pursue— or don’t— to address those concerns.

6. “The horse has left the barn”: Discursive closure and visual generative AI

Brause et al. (2023) have pointed out that media narratives and framings have a direct impact on policy making, economic investments, and research activities in emerging technological fields (p. 11), and Bareis and Katzenbach (2021) similarly argue that there is “a strong role for discourse in shaping the present and future sociotechnical pathways” of AI (p. 858). For these reasons, how the challenges and harms of visual generative AI are presented within the global news media have potentially immense consequences for how— or even if— these harms are remediated (see Shane, 2023).

The Lensa coverage suggests that, with few exceptions, there is a discourse of inevitability (Leonardi and Jackson, 2003; Markham, 2021) in play that fails to offer any novel or alternative framings— or even opportunities for discussion—to the harms posed by visual generative AI. The discourse about Magic Avatars certainly identifies those harms, and even strongly critiques them; however, it also treats them as intractable problems that are already too big to solve.

Although Lensa was arguably a novel technology at the time, it is remarkable how closely the press coverage about it hewed to existing framings and discourses about social media platforms and other types of algorithmic media. The discussions of predatory data practices closely align with popular discourses about users as exploited data subjects that took off exponentially in the wake of the Cambridge Analytical scandal (e.g., Isaac and Hanna, 2018; Koidl and Kapanova, 2020), the critiques of Lensa’s sexualized and racialized results map directly onto discourses about “algorithms of oppression” and other forms of gendered and racialized algorithmic biases (e.g., Noble, 2017; Benjamin, 2019), and analyses of Lensa user behaviors were plucked directly from the selfie moral panic playbook, with their insistence that

digital technologies amplify our most deleterious narcissistic tendencies, especially if “we” are women (Miltner and Baym, 2015; Senft and Baym, 2015; Abidin, 2016; Tiidenberg, 2018).

To be clear, algorithmic bias and predatory data practices are significant social concerns that are deserving of public attention (the gendered media panic around selfies and other self-representation practices, perhaps less so). However, it is not *that* these matters are being discussed but *how* they are being discussed that is the problem. By repeating existing framings of these issues, it reinscribes hegemonic understandings of them and perpetuates a sense of inevitability about the status quo. In other words, you might be upset that these technologies are stealing your data and making porn out of your photos, but there’s not much you can do about it. As Markham (2021) explains,

Focusing on how discourses are normalized or locked into repetitive loops helps specify how hegemony works in everyday practices. In systems of highly effective oppression or, what Gramsci labeled “control through consent,” people shut down alternatives themselves, naturalizing problems as “just the way things are.” (p. 392)

This framing of AI harms as “just the way things are” is present throughout the corpus, and particularly in the Predatory Data Practices and Biased Content coverage. This framing manifests in two key ways. The first is that AI simply reflects pre-existing societal biases: Prisma Labs spokesperson Anna Green explained to The New York Times that Lensa was not consciously applying biases, but rather that “essentially, A.I. is holding a mirror to our society” (Chen, 2023). This suggests that society— and not a series of corporate policy and design choices— is primarily responsible for these problems, an explanation that conveniently eschews the fact that generative AI systems have been proven to amplify and worsen societal biases (e.g., Nicoletti and Bass, 2023). The other framing tends to suggest that any potential interventions to address these concerns are simply too late. In describing the recourse available to people whose images are included in the LAION 5B dataset, TechCrunch’s Taylor Hatmaker (2022) explains that EU citizens can file takedown requests for individual images but “that’s about it”, because “the horse has already left the barn”. Shira Ovide (2022) of The Washington Post had a similarly fatalistic attitude about the creation of AI systems, stating that readers’ feelings about their data being used to train AI models were largely inconsequential, because they were powerless to stop it:

Good or bad, these AI systems are being built with pieces of you. What are the rules of the road now that you're breathing life into AI and can't imagine the outcomes? [...] Being part of the collective building of all these AI systems might feel unfair to you, or amazing. But it is happening.

Markham (2021) explains that inevitability and powerlessness can be outcomes of naturalization, which takes place when current elements of sociotechnical contexts are accepted simply as the way things are (p. 397). Even though Stable Diffusion and its ilk have been available to the public for less than two years, the idea that these models— and the systems they undergird— will inevitably cause harm seems to be already taken for granted. This is a phenomenon that Daniel Chandler (1995) refers to as the “technological imperative” (see also Widder et al., 2022). Chandler writes that “the technological imperative is a common assumption amongst commentators on 'new technologies'. They tell us, for instance, that the 'information technology revolution' is inevitably on its way and our task as users is to learn to cope with it.”

Indeed, the solution in the corpus suggested to those most harmed by visual generative AI was to simply abstain from using it if they didn’t like the images it was producing. However, given that AI is already incorporated into many key digital technologies, engaging in media refusal (Portwood-Stacer, 2013) is not always desirable or feasible, especially as it could result in exclusion from the public sphere for groups who are already excluded and marginalized, denying them important sources of community and resistance (e.g., Sobande et al., 2019). Furthermore, these kinds of logics place the onus on the user who is harmed, instead of on the technology (and corporations) that are responsible for the harm(s): if the solution for AI harms is for marginalized users is to refrain from using these systems, it offers a reprieve to those who own and control these technologies from having to address any underlying issues that they perpetuate.

In this way, we can see discursive closure at work. By acquiescing to a discourse of inevitability, the Lensa coverage offers “a particular view of reality that is maintained at the expense of equally plausible ones” (Deetz, 1992, p. 188). Consequently, any alternative solutions to the issues at hand end up being downplayed or ignored. In her Los Angeles Times op-ed, Molly Crabapple (2022) argued that “data sets such as LAION-5B must be deleted and rebuilt to consist only of voluntarily submitted work. AIs trained on copyrighted art must also be pulled.” This suggestion aligns with some recent work in computer science that recommends that “learning methods that physically manifest stereotypes or other harmful outcomes be paused, reworked, or even wound down when appropriate, until outcomes can be proven safe, effective, and just” (Hundt et al., 2022, p. 743). Some technology companies are already taking this approach: in February 2024, Google took Gemini, its newly-launched AI image generator, offline when it began producing offensive images (Griffin, 2024). And yet, these ostensibly reasonable ideas were found nowhere else in the press coverage of Lensa.

Markham (2021) explains that the power of anticipatory logics that flow through everyday discourse around technologies like AI end up building and reinforcing “a hegemonic ideology of external power and control” (p. 384). As a result, more expansive and/or creative approaches to the challenges and potential harms of emerging technologies end up being ignored or dismissed as unlikely, constraining the boundaries of possibility. Dandurand et al. (2023) suggest that the “indeterminacies and uncertainties” of novel technologies like visual generative AI can “become settled when a few experts close debates, obfuscate contrasting expectations, and shape the political economy of science and technologies” (p. 3). As the discourse about Lensa illustrates, this process is already happening in relation to visual generative AI.

7. Conclusion

This paper analyzed the global press coverage of visual generative AI app Lensa and found that the discourse surrounding it focused on the app’s predatory data practices, the biased content it produced, and the user behaviors associated with it. I argue that this coverage provides evidence of discursive closure around key issues associated with visual generative AI that supports the maintenance of the status quo. I suggest that the press coverage of Lensa, which both articulates key AI-related harms and frames those harms as intractable and insolvable, creates a discourse of inevitability that has implications for how these issues are understood by the public, and for the approaches that are taken to address them. In doing so, a more imaginative public discussion of what visual generative AI could– or should– be is foreclosed.

Markham (2021) argues that “people seem to have difficulty imagining futures in ways that do not reproduce current ideological trends or cede control and power to external, mostly corporate, stakeholders” (Markham, 2021, p. 385). Indeed, it can be difficult to imagine things differently if you are consistently provided with a similar set of limited visions over and over again; as Markham also points out, “this continuously repeated mantra of powerlessness to avoid what is inevitable is not just something we invent; it is learned (or taught) in micro doses through our news feeds” (p. 397). Part of the issue is how journalism– both a public service and a struggling industry– operates in the 21st century. As Ananny and Finn (2020) point out,

Today’s news happens not just where journalists are ready to see news happening, but where news infrastructures have been designed to look. Part of appreciating the myriad forces and values of the contemporary networked press means understanding how infrastructures of humans (journalists, designers, audiences) and nonhumans (data sets, algorithms, interfaces), together, construct ways of seeing social worlds and translating those visions into familiar forms of news (p. 1613)

The evolution of news into a “borderless global market” has shifted news production practices and media consumption habits, where major media outlets in Western countries are able to deliver content to international audiences at a low cost (Rafeeq and Jiang, 2018). Furthermore, the datafication of the

audience and the metrification of news coverage has transformed editorial and journalistic practices, with journalists choosing to write articles that they might not see as newsworthy but feel will perform well with audiences (Dodds et al., 2023). This is not to say that the issues raised in the Lensa coverage are not newsworthy; however, it might not be a huge leap to suggest that part of the reason that Lensa received so much attention in the press is because it was a topic that performed well, which contributed to widespread syndication and further coverage. It also might not be a huge leap to suggest that how Lensa was covered was likely influenced by previous framings or angles that had performed well in the past or would invariably do so; tech coverage is a beat for specific journalists and outlets, and it doesn't take a veteran editor to know that a story about a popular app making non-consensual pornography out of people's photos is certainly clickworthy.

However, while these journalistic circumstances and choices may be at least partial explanations as to why the Lensa discourse looks the way it does, it doesn't negate the problem posed by these determinist—and unimaginative—framings. As McKelvey et al. (2023) argued in the CFP for this Special Issue, AI scandals—like Lensa's generation of racialized, sexualized images combined with their exploitative data practices—may offer “easy, high-engagement stories” that are appealing to news outlets and audiences alike, but such news coverage also fails to offer meaningful public engagement, not to mention democratic praxis. Indeed, Deetz (1992) argues that “it should not be surprising that systems of domination are protected from careful exploration”. In the case of Lensa, it may seem counterintuitive that the extensive discussion of the downsides of visual generative AI would provide any distinct advantages to the corporate interests who dominate in this space, but in identifying these harms and then framing them as inexorable, the “careful exploration” of potentially inconvenient—and expensive—solutions is discarded. Identifying the problems of visual generative AI and then insisting they're unfixable is probably a more effective way to curtail meaningful discourse about these concerns than focusing exclusively on AI's benefits.

In the face of this, it may seem that altering the trajectory of the discourse about visual generative AI—not to mention its development and deployment—is unlikely or insurmountable. However, we need not yet wave the white flag; it is early days yet. The history of technology shows us that technological development—and how it is discussed—are by no means inevitable, and that things could always have been otherwise (e.g., Marx, 2010; Pinch and Bijker, 1984). While the challenges may be significant, change is not impossible with some imagination and a collective refusal to accept the status quo.

References

- Abidin, C. (2016). “Aren't These Just Young, Rich Women Doing Vain Things Online?": Influencer Selfies as Subversive Frivolity. *Social Media + Society*, 2(2), 2056305116641342. <https://doi.org/10.1177/2056305116641342>
- Adobe Communications Team. (2023, October 18). *Generation AI: Reimagining creativity with generative AI*. Adobe Blog. <https://blog.adobe.com/en/publish/2023/10/18/generation-ai-reimagining-creativity-with-generative-ai>
- Alder, R. S. (2022, December 17). The White Privilege And Racism Of Lensa's Magic Avatars. *ILLUMINATION-Curated*. <https://medium.com/illumination-curated/the-white-privilege-and-racism-of-lensas-magic-avatars-355b1b53d306>
- Alfano, M., Abedin, E., Reimann, R., Ferreira, M., & Cheong, M. (2024). Now you see me, now you don't: An exploration of religious exnomination in DALL-E. *Ethics and Information Technology*, 26(2), 27. <https://doi.org/10.1007/s10676-024-09760-y>
- Ananny, M., & Finn, M. (2020). Anticipatory news infrastructures: Seeing journalism's expectations of future publics in its sociotechnical systems. *New Media & Society*, 22(9), 1600–1618. <https://doi.org/10.1177/1461444820914873>
- Bareis, J., & Katzenbach, C. (2022). Talking AI into Being: The Narratives and Imaginaries of National AI Strategies and Their Performative Politics. *Science, Technology, & Human Values*, 47(5), 855–881. <https://doi.org/10.1177/01622439211030007>
- Beaumont, R. (2022, March 31). *LAION-5B: A NEW ERA OF OPEN LARGE-SCALE MULTI-MODAL DATASETS*. LAION. <https://laion.ai/blog/laion-5b>
- Benjamin, R. (2019). *Race after technology: Abolitionist tools for the new Jim Code*. John Wiley & Sons.
- Bianchi, F., Kalluri, P., Durmus, E., Ladhak, F., Cheng, M., Nozza, D., Hashimoto, T., Jurafsky, D., Zou, J., & Caliskan, A. (2023). Easily Accessible Text-to-Image Generation Amplifies Demographic Stereotypes at Large Scale. *Proceedings of*

- the 2023 ACM Conference on Fairness, Accountability, and Transparency, 1493–1504.
<https://doi.org/10.1145/3593013.3594095>
- Birhane, A., Prabhu, V. U., & Kahembwe, E. (2021). *Multimodal datasets: Misogyny, pornography, and malignant stereotypes* (arXiv:2110.01963). arXiv. <https://doi.org/10.48550/arXiv.2110.01963>
- Biron, B. (2023, January 16). *Sex, art theft, and privacy: Lensa exploded overnight — and now the avatar app is dealing with public backlash. Here's what you should know.* Business Insider. <https://www.businessinsider.com/lensa-ai-raises-serious-concerns-sexualization-art-theft-data-2023-1>
- Bou, C. P. (2023, February 12). *Assemblar-se a una ciborg per estar guapa: Com la intel·ligència artificial està modelant el cànon de bellesa.* *El Periódico*. <https://www.elperiodico.cat/ca/societat/20230212/bellesa-intel·ligencia-artificial-cara-ciborg-filtres-canons-82583384>
- Braithwaite, A. (2016). It's About Ethics in Games Journalism? Gamergaters and Geek Masculinity. *Social Media + Society*, 2(4), 2056305116672484. <https://doi.org/10.1177/2056305116672484>
- Braun, V., & Clarke, V. (2019). Reflecting on reflexive thematic analysis. *Qualitative Research in Sport, Exercise and Health*, 11(4), 589–597. <https://doi.org/10.1080/2159676X.2019.1628806>
- Brause, S. R., Zeng, J., Schäfer, M. S., & Katzenbach, C. (2023). Media representations of artificial intelligence: Surveying the field. In S. Lindgren (Ed.), *Handbook of Critical Studies of Artificial Intelligence* (pp. 277–288). Edward Elgar Publishing. <https://doi.org/10.4337/9781803928562.00030>
- Broussard, M., Diakopoulos, N., Guzman, A. L., Abebe, R., Dupagne, M., & Chuan, C.-H. (2019). Artificial Intelligence and Journalism. *Journalism & Mass Communication Quarterly*, 96(3), 673–695. <https://doi.org/10.1177/1077699019859901>
- Brown, D. (2023, July 13). The Viral AI App That's Triggering Baby Fever. *Wall Street Journal*. <https://www.wsj.com/articles/want-to-see-your-future-babies-this-viral-app-can-show-you-f4f65bdd>
- Buschek, C., & Thorp, J. (n.d.). *Models All The Way Down*. Knowing Machines. Retrieved 2 October 2024, from <https://knowingmachines.org/models-all-the-way>
- Byrne, D. (2022). A worked example of Braun and Clarke's approach to reflexive thematic analysis. *Quality & Quantity*, 56(3), 1391–1412. <https://doi.org/10.1007/s11135-021-01182-y>
- Celeb Surgeons Say Clients Asking to Look Like Lensa App AI Portraits*. (2022, December 12). TMZ. <https://www.tMZ.com/2022/12/12/plastic-surgeons-say-clients-asking-look-like-lensa-app-portraits/>
- Chandler, D. (1995). *Technological Determinism: The Technological Imperative*. Retrieved 20 May 2024, from <http://visual-memory.co.uk/daniel/Documents/tecdet/tcet07.html>
- Chen, B. X. (2022, December 21). How to Use ChatGPT and Still Be a Good Person. *The New York Times*. <https://www.nytimes.com/2022/12/21/technology/personaltech/how-to-use-chatgpt-ethically.html>
- Chuan, C.-H., Tsai, W.-H. S., & Cho, S. Y. (2019). Framing Artificial Intelligence in American Newspapers. *Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society*, 339–344. <https://doi.org/10.1145/3306618.3314285>
- Clark, N. (2022, December 20). Lensa's viral AI art creations were bound to hypersexualize users. *Polygon*. <https://www.polygon.com/23513386/ai-art-lensa-magic-avatars-artificial-intelligence-explained-stable-diffusion>
- Crabapple, M. (2022, December 21). Op-Ed: Beware a world where artists are replaced by robots. It's starting now. *Los Angeles Times*. <https://www.latimes.com/opinion/story/2022-12-21/artificial-intelligence-artists-stability-ai-digital-images>
- Dandurand, G., McKelvey, F., & Roberge, J. (2023). Freezing out: Legacy media's shaping of AI as a cold controversy. *Big Data & Society*, 10(2), 20539517231219242. <https://doi.org/10.1177/20539517231219242>
- Deetz, S. (1992). *Democracy in an age of corporate colonization: Developments in communication and the politics of everyday life*. SUNY press.
- Demopoulos, A. (2022, December 9). The inherent misogyny of AI portraits – Amelia Earhart rendered naked on a bed. *The Guardian*. <https://www.theguardian.com/us-news/2022/dec/09/lensa-ai-portraits-misogyny>
- Denton, E., Hanna, A., Amironesei, R., Smart, A., & Nicole, H. (2021). On the genealogy of machine learning datasets: A critical history of ImageNet. *Big Data & Society*, 8(2), 20539517211035955. <https://doi.org/10.1177/20539517211035955>
- Dihal K. (2024, May 11). Provocations & Sharing Concepts in AI & Culture [Session in *Global AI Cultures Workshop*]. ICLR 2024, Vienna, Austria [Online]. <https://iclr.cc/virtual/2024/workshop/20576>
- Dodds, T., De Vreese, C., Helberger, N., Resendez, V., & Seipp, T. (2023). Popularity-driven Metrics: Audience Analytics and Shifting Opinion Power to Digital Platforms. *Journalism Studies*, 24(3), 403–421. <https://doi.org/10.1080/1461670X.2023.2167104>
- Eapen, T. T., Finkenstadt, D. J., Folk, J., & Venkataswamy, L. (2023, July 1). How Generative AI Can Augment Human Creativity. *Harvard Business Review*. <https://hbr.org/2023/07/how-generative-ai-can-augment-human-creativity>
- Equity. (n.d.). *Stop AI Stealing the Show*. Equity. Retrieved 26 February 2024, from <https://www.equity.org.uk/campaigns-policy/stop-ai-stealing-the-show>
- Escalante-De Mattei, S. (2022, December 5). Careful—Lensa Is Using Your Photos to Train Their AI. *ARTnews*. <https://www.artnews.com/art-news/news/does-lensa-ai-use-your-face-data-for-selfies-1234649204/>
- Fast, E., & Horvitz, E. (2017). Long-Term Trends in the Public Perception of Artificial Intelligence. *Proceedings of the AAAI Conference on Artificial Intelligence*, 31(1), Article 1. <https://doi.org/10.1609/aaai.v31i1.10635>

- Flux, E. (2022, December 23). What does the rise of AI mean for the future of art? *The Age*. <https://www.theage.com.au/culture/art-and-design/what-does-the-rise-of-ai-mean-for-the-future-of-art-20221221-p5c81c.html>
- Gaskins, N. (2022, December 17). AI & Techno-Vernacular Creativity. *Medium*. <https://nettricegaskins.medium.com/ai-techno-vernacular-creativity-314281e1be68>
- Gillespie, T. (2024). Generative AI and the politics of visibility. *Big Data & Society*, 11(2), 20539517241252131. <https://doi.org/10.1177/20539517241252131>
- Griffin, A. (2024, February 22). Google takes AI image generator offline over racially diverse historical images. *The Independent*. <https://www.independent.co.uk/tech/google-gemini-ai-images-racially-diverse-b2500730.html>
- Haigney, S. (2023, January 11). AI's Best Trick Yet Is Showering Us With Attention. *The New York Times*. <https://www.nytimes.com/2023/01/11/magazine/ai-photo-filters-selfies.html>
- Hatmaker, T. (2022, December 5). Lensa AI, the app making 'magic avatars,' raises red flags for artists. *TechCrunch*. <https://techcrunch.com/2022/12/05/lensa-ai-app-store-magic-avatars-artists/>
- Heikkilä, M. (2022a, September 16). *This artist is dominating AI-generated art. And he's not happy about it*. MIT Technology Review. <https://www.technologyreview.com/2022/09/16/1059598/this-artist-is-dominating-ai-generated-art-and-hes-not-happy-about-it/>
- Heikkilä, M. (2022b, December 12). The viral AI avatar app Lensa undressed me—Without my consent. *MIT Technology Review*. <https://www.technologyreview.com/2022/12/12/1064751/the-viral-ai-avatar-app-lensa-undressed-me-without-my-consent/>
- Hill, L. W. (2022, June 12). I tried the AI photo app everyone was using. As someone with a history of disordered body image, it was very triggering to see the results. *Insider*. <https://www.insider.com/woman-with-disordered-body-image-tried-lensa-app-2022-12>
- Hundt, A., Agnew, W., Zeng, V., Kacianka, S., & Gombolay, M. (2022). Robots Enact Malignant Stereotypes. *2022 ACM Conference on Fairness, Accountability, and Transparency*, 743–756. <https://doi.org/10.1145/3531146.3533138>
- Hunter, T. (2023, April 14). AI selfies—And their critics—Are taking the internet by storm. *Washington Post*. <https://www.washingtonpost.com/technology/2022/12/08/lensa-ai-portraits/>
- Isaak, J., & Hanna, M. J. (2018). User Data Privacy: Facebook, Cambridge Analytica, and Privacy Protection. *Computer*, 51(8), 56–59. <https://doi.org/10.1109/MC.2018.3191268>
- Kamps, H. J. (2022, December 6). It's way too easy to trick Lensa AI into making NSFW images. *TechCrunch*. <https://techcrunch.com/2022/12/06/lensa-goes-nsfw/>
- Keller, S. M. (n.d.). AI Art Gives Users Gender Euphoria—But It's Not Without Controversy. *Out*. Retrieved 4 November 2023, from <https://www.out.com/art/2022/12/07/ai-art-gives-users-gender-euphoria-its-not-without-controversy>
- Kelly, C. (2022, December 11). Australian artists accuse popular AI imaging app of stealing content, call for stricter copyright laws. *The Guardian*. <https://www.theguardian.com/australia-news/2022/dec/12/australian-artists-accuse-popular-ai-imaging-app-of-stealing-content-call-for-stricter-copyright-laws>
- Khan, M. (2023, October 5). People are posting AI-generated yearbook pictures with this viral app. *CNBC*. <https://www.cnn.com/2023/10/05/how-to-create-ai-yearbook-pictures-with-viral-epik-app.html>
- King, M. (2022). Harmful biases in artificial intelligence. *The Lancet Psychiatry*, 9(11), e48. [https://doi.org/10.1016/S2215-0366\(22\)00312-1](https://doi.org/10.1016/S2215-0366(22)00312-1)
- Kircher, M. M., & Holtermann, C. (2022, December 8). How Is Everyone Making Those A.I. Selfies? *The New York Times*. <https://www.nytimes.com/2022/12/07/style/lensa-ai-selfies.html>
- Klee, M. (2022, December 12). A Psychologist Explains Why Your 'Hot AI Selfies' Might Make You Feel Worse. *Rolling Stone*. <https://www.rollingstone.com/culture/culture-features/lensa-app-hot-ai-selfie-self-esteem-1234644965>
- Koidl, K., & Kapanova, K. G. (2020). Introduction to the Special Issue: Re-Imagining a More Trustworthy Social Media Future. *Social Media + Society*, 6(2), 2056305120923007. <https://doi.org/10.1177/2056305120923007>
- Leonardi, P. M., & Jackson, M. H. (2004). Technological determinism and discursive closure in organizational mergers. *Journal of Organizational Change Management*, 17(6), 615–631. <https://doi.org/10.1108/09534810410564587>
- Lovejoy, B. (2023, February 15). AI chat may be big news, but it seems we're over AI photo apps. *9to5Mac*. <https://9to5mac.com/2023/02/15/ai-chat-ai-photo/>
- Luccioni, A. S., Akiki, C., Mitchell, M., & Jernite, Y. (2023). *Stable Bias: Analyzing Societal Representations in Diffusion Models* (arXiv:2303.11408). arXiv. <https://doi.org/10.48550/arXiv.2303.11408>
- Maimann, K. (2022, December 31). This image is raising money for a Toronto charity. The only problem? It's not real. *Toronto Star*. https://www.thestar.com/news/canada/this-image-is-raising-money-for-a-toronto-charity-the-only-problem-it-is-not/article_9ebc80eb-27ff-5e79-a6bd-acc44beb500f.html
- Marcus, J. (2022, December 10). Artists decry use of AI art: 'I'm concerned for the future of human creativity'. *The Independent*. <https://www.independent.co.uk/news/world/americas/ai-art-lensa-magic-avatar-b2242891.html>
- Markham, A. (2021). The limits of the imaginary: Challenges to intervening in future speculations of memory, data, and algorithms. *New Media & Society*, 23(2), 382–405. <https://doi.org/10.1177/1461444820929322>

- Marres, N. (2015). Why Map Issues? On Controversy Analysis as a Digital Method. *Science, Technology, & Human Values*, 40(5), 655–686. <https://doi.org/10.1177/0162243915574602>
- Martineau, K. (2021, February 9). *What is generative AI?* IBM Research Blog. <https://research.ibm.com/blog/what-is-generative-AI>
- Marx, L. (2010). Technology: The Emergence of a Hazardous Concept. *Technology and Culture*, 51(3), 561–577.
- McCluskey, M. (2022, December 14). Why It's So Hard to Resist Turning Your Selfies Into Lensa AI Art. *Time*. <https://time.com/6240648/lensa-ai-psychology-behind/>
- McKelvey, F., Redden, J., Roberge, J., and Stark, L. (2023, May 24). *CFP: Special issue on (un)Stable Diffusions - Machine Agencies*. <https://machineagencies.milieux.ca/jdsr-cfp/>
- Meyerowitz, A. (2022, December 9). Megan Fox complains that her AI profile photos have been sexualised by viral app. *Glamour UK*. <https://www.glamourmagazine.co.uk/article/megan-fox-lensa>
- Miltner, K. M., & Baym, N. K. (2015). The Selfie of the Year of the Selfie: Reflections on a Media Scandal. *International Journal of Communication*, 9, 1701–1715.
- Miltner, K. M., & Highfield, T. (2024). *The Possibilities of 'Good' Generative AI in the Cultural and Creative Industries* (p. 10). The British Academy. <https://www.thebritishacademy.ac.uk/publications/the-possibilities-of-good-generative-ai-in-the-cultural-and-creative-industries/>
- Natale, S., & Ballatore, A. (2020). Imagining the thinking machine: Technological myths and the rise of artificial intelligence. *Convergence*, 26(1), 3–18. <https://doi.org/10.1177/1354856517715164>
- Nicoletti, L., & Bass, D. (2023, June 7). Humans Are Biased. Generative AI Is Even Worse. *Bloomberg*. <https://www.bloomberg.com/graphics/2023-generative-ai-bias/>
- Noble, S. U. (2018). *Algorithms of oppression: How search engines reinforce racism*. NYU Press.
- O'Brien, S. A. (2022, December 11). 'Magic' AI Avatars Are Already Losing Their Charm. *Wall Street Journal*. <https://www.wsj.com/articles/lensa-ai-avatars-women-11670705062>
- Offert, F., & Phan, T. (2022). *A Sign That Spells: DALL-E 2, Invisual Images and The Racial Politics of Feature Space* (arXiv:2211.06323). arXiv. <https://doi.org/10.48550/arXiv.2211.06323>
- Ouchchy, L., Coin, A., & Dubljević, V. (2020). AI in the headlines: The portrayal of the ethical issues of artificial intelligence in the media. *AI & SOCIETY*, 35(4), 927–936. <https://doi.org/10.1007/s00146-020-00965-5>
- Ovide, S. (2022, December 10). Your selfies are helping AI learn. You did not consent to this. *Washington Post*. <https://www.washingtonpost.com/technology/2022/12/09/chatgpt-lensa-ai-ethics/>
- Perez, S. (2023, February 13). The AI photo app trend has already fizzled, new data shows. *TechCrunch*. <https://techcrunch.com/2023/02/13/the-ai-photo-app-trend-has-already-fizzled-new-data-shows/>
- Pinch, T. J., & Bijker, W. E. (1984). The Social Construction of Facts and Artefacts: Or How the Sociology of Science and the Sociology of Technology might Benefit Each Other. *Social Studies of Science*, 14(3), 399–441. <https://doi.org/10.1177/030631284014003004>
- Portwood-Stacer, L. (2013). Media refusal and conspicuous non-consumption: The performative and political dimensions of Facebook abstention. *New Media & Society*, 15(7), 1041–1057. <https://doi.org/10.1177/1461444812465139>
- Prisma Labs, Inc. (n.d.). *Lensa's Magic Avatars Explained*. Notion. Retrieved 22 November 2023, from <https://prismalabs.notion.site/prismalabs/Lensa-s-Magic-Avatars-Explained-c08c3c34f75a42518b8621cc89fd3d3f>
- Rafeeq, A., & Jiang, S. (2018). From the Big Three to elite news sources: A shift in international news flow in three online newspapers TheNational.ae, Nst.com.my, and Nzherald.co.nz. *The Journal of International Communication*, 24(1), 96–114. <https://doi.org/10.1080/13216597.2018.1444663>
- Rekhi, D. (2022, December 17). *Cyber experts train lens on Lensa-like apps—ET Telecom*. ETTelecom.Com. <https://telecom.economictimes.indiatimes.com/news/cyber-experts-train-lens-on-lensa-like-apps/96292115>
- Roberts, M. (2022, May 12). The Euphoric Highs & Problematic Lows Of AI Avatar Art. *Refinery20*=9. <https://www.refinery29.com/en-gb/ai-avatar-art>
- Roose, K. (2022, October 21). A.I.-Generated Art Is Already Transforming Creative Work. *The New York Times*. <https://www.nytimes.com/2022/10/21/technology/ai-generated-art-jobs-dall-e-2.html>
- Sandoval-Martin, T., & Martínez-Sanzo, E. (2024). Perpetuation of Gender Bias in Visual Representation of Professions in the Generative AI Tools DALL·E and Bing Image Creator. *Social Sciences*, 13(5), Article 5. <https://doi.org/10.3390/socsci13050250>
- Santos, K. M. L. (2020). The bitches of Boys Love comics: The pornographic response of Japan's rotten women. *Porn Studies*, 7(3), 279–290. <https://doi.org/10.1080/23268743.2020.1726204>
- Scherer, M. (2024, January 4). The SAG-AFTRA Strike is Over, But the AI Fight in Hollywood is Just Beginning. *Center for Democracy and Technology*. <https://cdt.org/insights/the-sag-aftra-strike-is-over-but-the-ai-fight-in-hollywood-is-just-beginning/>
- Scheufele, D. A., & Lewenstein, B. V. (2005). The Public and Nanotechnology: How Citizens Make Sense of Emerging Technologies. *Journal of Nanoparticle Research*, 7(6), 659–667. <https://doi.org/10.1007/s11051-005-7526-2>
- Senft, T. M., & Baym, N. K. (2015). Selfies Introduction ~ What Does the Selfie Say? Investigating a Global Phenomenon. *International Journal of Communication*, 9(0), Article 0.

- Shane, T.S. (2023). AI incidents and ‘networked trouble’: The case for a research agenda. *Big Data & Society*, 10(2), 20539517231215360. <https://doi.org/10.1177/20539517231215360>
- Silberling, A. (2022, December 1). Lensa AI climbs the App Store charts as its ‘magic avatars’ go viral. *TechCrunch*. <https://techcrunch.com/2022/12/01/lensa-ai-climbs-the-app-store-charts-as-its-magic-avatars-go-viral/>
- Smith, A. (2022, December 20). Lensa AI image app helps some trans people to embrace themselves | Context. *Openly News*. <https://www.openlynews.com/i/?id=0d0b52f6-0fa4-4c27-9881-f689ad9fbe87>
- Snow, O. (n.d.). ‘Magic Avatar’ App Lensa Generated Nudes From My Childhood Photos. *WIRED*. Retrieved 29 November 2023, from <https://www.WIRED.com/story/lensa-artificial-intelligence-csem/>
- Sobande, F., Fearfull, A., & Brownlie, D. (2020). Resisting media marginalisation: Black women’s digital content and collectivity. *Consumption Markets & Culture*, 23(5), 413–428. <https://doi.org/10.1080/10253866.2019.1571491>
- Stability AI. (2022, August 22). *Stable Diffusion Public Release*. Stability AI. <https://stability.ai/news/stable-diffusion-public-release>
- Sun, M., & Harmon, S. (2024, March 6). ‘The worst AI-generated artwork we’ve seen’: Queensland Symphony Orchestra’s Facebook ad fail. *The Guardian*. <https://www.theguardian.com/culture/2024/mar/06/queensland-symphony-orchestra-ai-facebook-ad-criticism>
- Sung, M. (2022, December 7). Lensa, the AI portrait app, has soared in popularity. But many artists question the ethics of AI art. *NBC News*. <https://www.nbcnews.com/tech/internet/lensa-ai-artist-controversy-ethics-privacy-rcna60242>
- Thiel, D. (2023, December 20). Investigation Finds AI Image Generation Models Trained on Child Abuse. *Stanford Cyber Policy Center*. <https://cyber.fsi.stanford.edu/news/investigation-finds-ai-image-generation-models-trained-child-abuse>
- Tiidenberg, K. (2018). *Selfies: Why We Love (and Hate) Them*. Emerald Publishing.
- Tran, T. H. (2022, December 10). Image Apps Like Lensa AI Are Sweeping the Internet—And Stealing From Artists. *The Daily Beast*. <https://www.thedailybeast.com/how-lensa-ai-and-image-generators-steal-from-artists>
- Westerlund, M. (2019). The Emergence of Deepfake Technology: A Review. *Technology Innovation Management Review*, 9(11), 40–53. <https://doi.org/10.22215/timreview/1282>
- Widder, D. G., Nafus, D., Dabbish, L., & Herbsleb, J. (2022). Limits and Possibilities for “Ethical AI” in Open Source: A Study of Deepfakes. *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency*, 2035–2046. <https://doi.org/10.1145/3531146.3533779>