

**My more-than-human digital twin: embodiment, feminist AI,
and the struggle for representation**

PROPHET, Jane

Available from Sheffield Hallam University Research Archive (SHURA) at:

<https://shura.shu.ac.uk/36302/>

This document is the Published Version [VoR]

Citation:

PROPHET, Jane (2025). My more-than-human digital twin: embodiment, feminist AI,
and the struggle for representation. AI & SOCIETY. [Article]

Copyright and re-use policy

See <http://shura.shu.ac.uk/information.html>



My more-than-human digital twin: embodiment, feminist AI, and the struggle for representation

Jane Prophet¹

Received: 17 March 2025 / Accepted: 24 September 2025
© The Author(s) 2025

Abstract

Artificial intelligence is an entangled, more-than-human relational network, shaping and being shaped by the societal, cultural, and political structures in which it is embedded. This paper explores the role of artists as critical practitioners engaging with AI, to examine how AI-generated self-representations materialise identity and reinforce or counter well-known AI biases in gender, race, and embodiment. Drawing on feminist technoscience – particularly its focus on the entanglement of body, environment, and technology – and autotheory, the study treats generative AI as an instrument of vision and voice. Using generative AI to create a partial digital twin, a self-portrait, is situated as an inherently embodied and relational practice. Through a combination of desk-based research and a practice-based, reflexive engagement with RunwayML, the paper documents the author’s attempt to create an identical-looking digital twin, revealing systemic biases embedded in AI-generated self-portraits. The paper uses embodiment as connective tissue linking theoretical and lived experience. Generative AI consistently misrepresented gender and age, defaulting to hyper-feminized aesthetics and youthful features while reliably reproducing Whiteness. The study also critically examines voice cloning and text-to-speech synthesis, highlighting how AI’s training data constrain language, accent, and vocal traits. By positioning AI-generated imagery and voice synthesis as material-discursive practices, the research extends debates on bias, agency, and self-representation in human–machine interactions. It argues that artists working with generative AI not only expose its epistemic limitations but also provide counter-narratives through creative, embodied interventions. The findings highlight the ways artists can help to meet the urgent need for more inclusive AI infrastructures, transparent dataset practices, and a reframing of digital self-representation beyond generative AI’s algorithmic defaults.

Keywords Feminism · Gender · Embodiment · AI image generation · Prompting · Autotheory · Videoart

1 Introduction

Digital twins, typically rooted in industrial and biomedical domains, function as dynamic simulations using real-time sensor data and analytical processing to mirror physical entities (Jones et al 2020 n.p.). In contrast, this research focuses on simpler AI-generated video, cloned voice, and image-based self-portraits—artistic and expressive forms that foreground identity and embodiment. While distinct in application, digital twins and the project described both employ machine learning (ML) trained on personal data to construct a digital proxy of the self, challenging conventional

boundaries of subjectivity and raising critical questions about presence, agency, and authenticity in digital self-representation. The author takes a practice-based research approach, creating an AI-generated video self-portrait to explore how race, gender, and identity are algorithmically shaped, constrained, or erased and how that can be countered through producing AI-generated self-portraits, a precursor to developing a digital twin that would necessitate incorporating real-time biometric, location, and activity data. The self-portraits do not share digital twins’ characteristics of interacting in real-time and bridging physical and virtual systems. They are closer to avatars, virtual representations of a person that mimic their physical appearance, movements, voice, and expressions. This discussion addresses a subset of the special issue’s theme, Digital Identity & Self-Representation: Self-portrait, digital footprint, digital identity, and personae management.

✉ Jane Prophet
j.prophet@shu.ac.uk

¹ Sheffield Creative Industries Institute, Sheffield Hallam University, Sheffield, UK

The self-portraits are part of a video artwork with a self-reflexive narrative that describes the artist/author using an AI image generator, RunwayML, to create an audio-visual, identical-looking digital twin. The goal was to produce realistic-looking ‘representational images’, accurate video depictions of the 60 year-old artist in real-world settings. However, creating images of an AI character recognisable as the author was a struggle. Most images produced with RunwayML using text and image prompting were not recognisable as self-portraits because the character had so many characteristics that the author did not share. However, some were close enough to be uncanny. It became clear these were ‘presentational images’. Self-presentational images show us performing differently – presenting ourselves differently – depending on our context. Seeing a version of herself, in the form of a digital twin who was unfamiliar, in particular because of the presentation of gender and age, was alienating to the author. This paper asks, “How do AI-generated self-representations materialise identity?” “What does the process of co-creating with AI reveal about embodiment, bias, and agency in human–machine interactions?” “How do our identities and perceptions of ourselves shift when we see and hear AI-generated representations of ourselves?”.

Donna Haraway’s assertion that all scientific observation is partial, situated (Haraway 2013, p. 582), and mediated aligns with Karen Barad’s concept of “mattering” as the entangled emergence of meaning and materiality (Barad 2007, p. 70). These frameworks shaped the author’s methodological approach, which combined desk-based research and creative practice using RunwayML to explore how bodies are differentially rendered intelligible and consequential within generative AI’s sociotechnical systems. The literature review opens with key artworks, then turns to relevant text-based scholarship. Instead of treating generative AI as an independent tool, feminist technoscience theories help to position it as part of a complex ‘always becoming’ more-than-human relational network. AI materialised through datasets, algorithms, and generated representations, is not neutral – it expresses the perspectives of those bodies that are made legible, valued, and included in these digital ontologies, and not of those rendered invisible or misrepresented. Revisiting feminist and critical race theories of objectification helps reveal processes of mattering or materialising AI by counting bodies that ‘count’ and revealing quantitative and qualitative differences between human bodies in AI’s more-than-human relational network. This moves from a discussion about images of human bodies to how the body sounds when speaking, connecting alt-text, text, language and voice.

The methodology section builds a scaffold, an inseparable connective tissue —a fascia — that connects literature to method and supports the author’s later discussion of their embodied engagement with generative AI. The scaffolding

combines autoethnography, reflexive practice and autotheory. Autoethnography supports an exploration of textuality within AI instruments, from invisible alt-texts embedded in commonly used training datasets to writing prompts and voice-over scripts as a practice of writing the self. The fieldwork – creating a video with generative AI – uses a reflexive practice-based research process that depends on iteratively writing, observing, and listening to what emerges in collaboration with AI. As a hybrid methodology, autotheory enabled the artist/author to merge personal experience with critical theory, positioning their lived, embodied experience of AI video generation as a site of knowledge production. By combining these methods, the research critically engages with the ‘black box’ of AI, interrogating how AI generates images and text and materialises identity, embodiment, and agency.

2 Literature review

2.1 Key artworks in the field head

Artificial intelligence is co-constituted as part of an entangled more-than-human relational network, continuously shaping and being shaped by societies, cultures, and politics with which it emerges. As AI-generated images, voices, and narratives become increasingly integrated into artistic and creative practices, artists working with AI offer critical, embodied perspectives on AI biases, constraints, and epistemologies. Their engagements have ripple effects on the relational network. Since the mid-1990s, Hannah Redler-Hawes (Redler-Hawes n.d. n.p.) has been curating visual artists who work with data and, increasingly, engage with AI. My collaborative artworks are a small part of this trajectory, from my material-discursive explorations of the World Wide Web through creating real-time online artificial life artworks in the 1990s (Prophet 2001), to artworks created through collaborative bioinformatic research and simulations based on formal models of stem cell behaviour (d’Inverno 2006), to generating visually convincing tree forms that express air quality data in real-time (Prophet et al. 2018). One influence is the collaborative duo Tessa Elliott and Jonathan Jones Morris. Redler-Hawes commissioned their 1999 “interadditive” artwork *Machination* (Arnold et al. 2024 n.p.) which included a self-learning neural network that referred to a hand-crafted database of 1000 household objects. *Machination* shows the importance of creating “counter datasets” that challenge dominant narratives and offer alternative, more inclusive forms of representation. This is especially pressing now and a driver for my approach to training AI-image generators with images I have created from both 35mm slides of plant materials and photographs of myself. Anna Ridler’s AI-generated work, *Mosaic Virus* (Ridler 2019 n.p.), is a

contemporary exploration of capitalism and collapse that refers to tulip and bitcoin mania. Ridler's custom database of 10,000 photographs of tulips, which has increased AI literacy (Hemment et al. 2023 p.6) is a reflection on how AI-generated images are not an image of a real tulip, but what AI thinks a tulip should be. To produce *herbAlrium* (Prophet 2024 n.p.) I created numerous smaller-scale ML datasets and sent consistent feedback to Runway when prompts did not result in salient images, to influence the underlying AI literacy. For *Prosthetic Memory* artist-technologist M Eifler makes a bespoke database of video clips recorded for their future self, which are displayed when a previous date or keyword is spoken. The work partially replaces Eifler's long-term memory, which was damaged by a childhood brain injury. Eifler explores questions central to my work with generative AI, asking, "Do our assumptions, fears, and uses for AI change when data and ML models are created by individuals and families at personal, instead of corporate, scale?" (Eifler, 2020 n.p.) I take this up through *herbAlrium*'s narrative form, an illustrated critique of corporate-scale AI and ML models. Eifler also asks, "Given our increasing exposure to algorithmic interventions, how do our identities and perceptions shift when we see ourselves and others through that lens?" While not the impetus at the start of my practice-based research, this became a key secondary question as I began using AI to generate image and voice self-portraits. How perceptions of ourselves alter when we hear our cloned voice was also salient when I created an AI-generated voice. Stephanie Dinkins has critically engaged with this in her chatbot-making practice, as expressed in the ongoing project *Not The Only One* (N'TOO), actively intervening to counter erasures of Black voices in existing chatbot datasets, constructing new AI text and voice. Dinkins writes "[N'TOO] is a voice-interactive AI entity designed, trained, and aligned with concerns and ideals of people underrepresented in the tech sector" (Dinkins n.p.).

Strategies for critical artistic engagements with AI are not limited to producing new curated datasets. Trevor Paglan's work interrogates AI's "invisible infrastructures"; his work *ImageNet Roulette* draws "attention to the things that can – and regularly do – go wrong when artificial intelligence models are trained on problematic training data" (Crawford and Paglan n.d n.p.). My video, *herbAlrium* (Prophet 2024 n.p.), in combination with this paper, similarly draws attention to what goes wrong when we work with invisible infrastructures, with a particular focus on gender and age bias. Feminist and intersectional approaches to AI include those taken by visual artist Minne Atairu, who used Midjourney's (V4) text-to-image algorithm to generate her *Blond Braids Study*, "studio portraits of 'blue-black' or 'plum-black' complexioned twins sporting blonde braids" (Atairu 2023 n.p.). Atairu notes that this unnaturally silky, blond hair exposed racial biases

in AI-generated imagery, showing how training datasets fail to account for diverse forms of Black representation. I playfully use the example of AI-generated images of hands, presented with perfect manicures and wedding rings whenever the prompt includes "woman" to flag up the AI-generated images' lack of diverse femme presentations. Christine Liao uses DALL-E 3's AI text-to-image generator to create their avatar Liliann "to challenge stereotypical, hypersexualized, heteronormative White Eurocentric standard beauty and representations of femininity in avatars" (Liao 2024 p.12). Liao argues that a critical and creative engagement with AI avatar-making is vital to counter questionable gender representations amplified by these avatar images. Later, I describe the frustration of prompting generative AI to present me with representations that are more androgenous, closer to my real-life appearance. Liao generates still images using art methods such as image-based exquisite corpse, glitch, and remix. These examples demonstrate that artists' co-creation with AI is not merely an aesthetic exploration but a crucial site of critique, resistance, and knowledge production. By foregrounding embodiment and intersectionality in a discussion of creating a digital twin self-portrait, this paper argues that artistic practice offers a material-discursive engagement with AI.

While existing research explores prompt engineering, AI aesthetics, and dataset bias, there is limited scholarship, with Liao's notable exception, on how artists and those with expertise in self-portraiture experience and navigate AI's constraints when using AI-generated images and videos to create digital twins or self-portraits. Computational social scientist Luhang Sun and co-authors studied presentational biases in AI, calling for more research "to examine the potential effects of exposure to gendered AI-generated images, and explore strategies to effectively mitigate gender biases in AI Models." (Sun et al. 2024 p.13). The collaboration between the author/artist and AI, somewhat inadvertently, responded to that call. The author's curiosity about how 'failed' and alienating self-portraits were instantiated led to more rigorous prompting work with RunwayML, with observational notes recorded along the way. Concurrent desk-based research revealed strategies to counter stereotypical representations. This iterative process of generative prompting and critical reflection highlighted not only the limitations of AI image-making but also the embodied labour required to resist its defaults. As the author navigated between algorithmic constraints and representational agency, embodiment emerged as a central concern—shaping both the creative process and the theoretical lens. Embodiment is the connective tissue that holds together different theories, lived experiences, and technologies (like AI) as they interact in fluid, unpredictable, and context-dependent ways.

2.2 Situating AI instruments of vision

In her theory of situated knowledge, Donna Haraway uses the metaphor of vision to critique objectivity and knowledge production in science. Her assertion that vision in relation to scientific observation is always partial, located, and mediated by instruments of vision – microscopes, MRI scanners, and imaging technologies – is a key element of this discussion of AI image generators, which are treated as instruments of vision. She argues, “[v]ision requires instruments of vision; an optics is a politics of positioning” (Haraway 2013, p. 582). Instruments of vision imply an embodied observer, and Haraway emphasises the importance of the human body more broadly – showing that situated knowledge is produced through lived, bodily experiences, vision is embodied and positioned – there is no neutral or disembodied observer. This lack of neutrality applies to the technologies, environments, and power structures that generate knowledge, in this case, AI instruments of vision and how they are created and used.

The AI instrument of vision known as the Large-scale Artificial Intelligence Open Network (LAION) dataset is a huge open-source collection of image and alt-text pairs used to train numerous AI systems. AI Image generation models like MidJourney and Stable Diffusion are known to utilise LAION, as do models such as RunwayML, built on Stable Diffusion. As a result, co-creating self-portraits with generative AI is affected by any bias in the LAION datasets. LAION’s relational network includes embodied humans, each with their own positionality, who created billions of images and their paired alt-texts which are stored online; human programmers who write algorithms that scrape the internet to gather those images and texts into datasets; human ML engineers who use those datasets to train AI; more humans, like me, working from our homes and offices, using those AI image generators to generate images that we then, in turn share.

Pragmatic approaches such as reviewing literature that exposes more about various entangled parts of generative AI’s relational network (Hancox-Li And Kumar 2021, n.p.) can help us to grasp the complexity of these vast relational networks. We can decipher how this knowledge is situated by considering the images and alt-texts that form the datasets the ML is trained on, and by asking a series of questions that address embodiment: *who* created them, *which* human bodies are made (in) visible through those images and texts, and *how* are those bodies depicted? Situated knowledge requires researchers to be accountable and acknowledge their positionality, as well as the ethical implications of their work. People developing AI generators typically do not disclose their positionality. Still, their work and work product have been analysed by numerous researchers in attempts to situate AI knowledge more clearly. Learning more about embodied

human contexts in which datasets are produced helps reveal “visualizing tricks and powers of modern sciences and technologies that have transformed the objectivity debates” (Haraway 2013, p. 582). Firstly, what do we know about the human bodies depicted in the billions of pairs of images and alt-text descriptions that have been scraped from the internet and used in datasets selected by engineers for ML?

2.3 Objectification in image and alt-text pairs

Objectification is central to feminist and critical race theories. It refers to treating a person as an object rather than a subject with feelings and experiences, typically reducing the body to a passive object for visual consumption—stripping away agency, context, and complexity. Objectification theory (Fredrickson and Roberts, 1997) argues that women are routinely treated as bodies or body parts, valued mainly for their use to others. Objectification is also intersectional (Crenshaw 1989, p.140), shaped by the interplay of race, gender, sexual orientation, class, ability, and other categories, with overlapping identities intensifying objectification and oppression. Objectification, though not always intersectionality, underpins feminist analyses of beauty, showing how women internalise an outsider’s gaze, leading to self-surveillance and disempowerment. Laura Mulvey’s theory of the “male gaze” (Mulvey, 1989, p.19), where women are presented as passive sexual objects for the pleasure of the heterosexual male viewer, remains particularly insightful for analysing images in generative AI training datasets. The assumed gaze is also White, reflecting dominant beauty standards that privilege Whiteness, thinness, and Eurocentric features, as critiqued by Minne Atairu in *Blond Braids Study*. This results in racialised objectification, where women of colour are hypersexualised or desexualised in contrast to White women. Psychological studies show that internalising the objectifying White male gaze increases women’s body shame and self-presentation anxiety, turning us into visual objects (Calogero, 2004, p.20).

Understanding how objectification operates in generative AI helps us to interrupt it, though not easily. Even when people document themselves to resist objectification, their images are often appropriated and repurposed in ways that reinforce racist and sexist norms. Francesca Sobande’s work on digital racism reveals how representations of Black people are shaped by racialised digital marketplace logics and (re)mediations of Blackness. Her research on computer-generated racialised influencers—lifelike digital creations with lucrative online profiles—highlights how brands engage CGI micro-celebrities that conform to stereotypical beauty ideals. Sobande notes that White men have created some of the most commercially successful CGI Black feminine-presenting influencers to reflect White heterosexual male ideas of Black femininity (Sobande 2021, p.135).

In image-alt-text databases, objectification can be incorporated into images, and/or their associated alt-texts. Binary divisions of gender, such as images with alt-text labels like “man” or “woman,” reinforce heteronormative perspectives. From now on, I will intentionally use more gender-inclusive terms – “feminine-presenting” and “masculine-presenting” – although these terms are rarely used in the datasets discussed. For clarity, I will repeat binary terms used in datasets and by other writers but put them in quotation marks as a reminder of their potential partiality.

It is estimated that less than 1% of images on the internet currently include alt-text, meaning 99% of images would be excluded from LAION’s dataset. Alt-text is not simply a description of what is in photos – it guides the human gaze, prompts a particular interpretation of visual information, and acts as an instruction on how to look. Who tags images with alt-text, and what meanings do those texts assign for human users and bots? Demographic data about the humans (and the AI) who created the billions of pictures and alt-texts used in LAION is not available at any granularity. Still, researchers who have studied LAION data, looked closely at the images and read associated alt-texts remind us that oppressive sexualised and racialised meanings are common, even when photos themselves are benign. Partly because of where data is scraped from, alt-text-image pairings are often not congruent:

“the alt text associated with such [pornographic] images, which may have a relative benign representation in the purely textual context, is often perverted through the lens for sociocultural fetishizations of the same terms in the visual context. For example, in the LAION-400 M dataset, words such as ‘mom’, ‘nun’, ‘sister’, ‘daughter’, ‘daddy’ and ‘mother’ appear with high frequency in alt text for sexually explicit content. We have also observed a similar effect in the reverse direction, e.g. where innocent images of school girls have alt-text that is loaded with terms typically searched for by paedophiles and sexual predators.” (Birhane et al. 2021 n.p.).

2.4 Externalized appearances and normalisation in datasets

From the example above, it is easy to assume that there are more images of feminine-presenting bodies than masculine-presenting bodies on these datasets. Further research indicates that female-identified images are underrepresented by 10–15% (Sun et al. 2024, p.6). Male over-representation may hinder creative practitioners seeking to generate female-identified self-images. Moreover, female bodies are more frequently rendered in sexualized poses, shaped by the pairing of images with sexualized alt-text. Cognitive scientists analysing the LAION-400 M dataset found “troublesome and explicit images and text pairs of rape,

pornography, malign stereotypes, racial and ethnic slurs, and other extremely problematic content” (Birhane et al. 2021, n.p.). Their audit also revealed weak and misleading links between image content and its textual description—for instance, Safe For Work (SFW) images were often paired with Not Safe For Work (NSFW) alt-text. Further research shows that AI-generated alt-texts and image captions frequently perpetuate sexual and racial objectification, often ignoring emotional nuance—especially for partially clothed subjects. Intersectional critiques remain limited, particularly around age. Some work notes the sexualization of images depicting children, and that harmless photos of children are frequently accompanied by sexualised alt-text; systematic analyses of older individuals are lacking. Studies on facial ageing show that age-related changes in expression lead to misinterpretations—such as reading neutral expressions as sadness or anger. AI is likely to describe images of older, unsmiling female-identified subjects using terms like ‘angry’ or ‘sad’.

In their finely woven paper, *How We’ve Taught Algorithms to See Identity: Constructing Race and Gender in Image Databases for Facial Analysis*, Scheuerman et al. (2020 p.8) critique how facial analysis datasets treat race and gender as externally legible, singular categories. Their interdisciplinary team notes that the annotations accompanying images often conflate gender presentation with sex and interpret race based on phenotypic features. These practices ignore the social construction of identity and contribute to representational narrowing, especially when these datasets inform AI image generation (Buolamwini And Gebru 2018, p.4). By analysing whose bodies appear— and how machines and humans are guided to categorise them via annotations and alt-texts—we observe how binary classifications and objectification constrain variation.

Little is known about the positionality of ML engineers who train these systems. Demographic data shows that three-quarters identify as men (Young et al. 2021, n.p.), and less than a quarter come from minoritised racial or ethnic groups. These statistics invite scrutiny of claims that ML methods are universally objective (Hancox-Li And Kumar 2021, n.p.). Far from a neutral “view from nowhere,” image and alt-text training data often reflect dominant perspectives—filtering AI’s gaze through norms that are human, male, White, and heterosexual, or, when unfiltered, reflect broader patriarchal, racist, and capitalist logics.

This section situates specific human contributors within AI’s more-than-human relational network. Understanding their identities—which cannot be conflated with what we know about their bodies—helps trace how biases are embedded in the development of AI instruments of vision. Ocular-centric paradigms across fields like art and science privilege visuality, drawing on metaphors such as reflection, refraction and the gaze. This analysis counters ocularcentrism

by drawing on Haraway's notion of embodied vision and expanding content analysis of training data to include alt-text. In summary, meaning in these systems emerges not just from images, but from their pairing with language. Alt-text, often treated by researchers without visual impairments as a proxy for sight, is designed to be spoken. It assumes a voice. While some people with low vision prefer to engage visually, alt-text offers aural access to visual content. However, its rhetorical and sensory dimensions are often overlooked in discussions of accessibility. Its dual role—as descriptive metadata and as speech—deserves more critical attention in understanding how AI mediates vision and voice.

3 AI instruments of (whose) embodied voices?

Materialising identity through video self-representations often entails more than creating images; it gives these self-representations, or self-portraits, a voice, a process that results in AI-generated characters that are more like digital twins. Scholars like Mara Mills (Mills 2012 p.36) and Jonathan Sterne discuss technologies of sound, such as telephony and hearing aids, as 'instruments of listening' that shape what and how we hear, much like Haraway's instruments of vision influence what and how we see, with Sterne pointing towards situatedness, noting "Hearing requires positionality" (Sterne 2012 p.4). Independent researchers invoke the link between text and speech in their audit of LAION, which they describe as a "vision-language dataset" (Birhane et al. 2023 n.p.). This critical shift from 'alt-text', from text as writing or reading, to 'language' invokes human voice. The audit determined that harmful text-based content is not filtered out consistently; therefore, datasets incorporate misogynistic and racist hate speech, much as they are co-constituted with misogynistic and racist images. They show that Not Safe For Work (NSFW) labelling is somewhat successful in filtering image data and removing objectifying photos, some of which may be paired with hate speech, which is also removed. However, hate speech in text paired with benign images is less likely to be tracked and filtered out.

In *Print Is Flat, Code Is Deep* N. Katherine Hayles highlights that the medium through which texts are instantiated matters – materiality is an emergent property, that cannot be specified in advance, but is "open to debate and interpretation, ensuring that discussions about text's "meaning" will also take into account its physical specificity as well." (Hayles 2004 p.67) Screen readers and other text-to-speech (TTS) instruments, including AI voice-over instruments, materialise text as language spoken by AI-generated voices; they are intermediaries that filter and translate text's content, inflecting it in ways that can reinforce or counter power structures. Each human voice is unique, altering uniquely

in response to changes in the human body. For example, testosterone causes the elongation and thickening of vocal folds, leading to changes in the voice. Human voices also change in response to the environment, illness and ageing and are shaped by other human voices they hear. Traces of how we have been situated differently geographically and in terms of class are just two examples of how our identity is materialised through our unique voices. In summary, our voices and speech patterns are nuanced, and speech that does not match the listener's expected pronunciation patterns can cause listener fatigue, misinterpretation, and lead to disengagement or a lack of trust.

Most existing AI-generated voices lack diversity and nuance. Understanding how these voices are trained shows why this is the case. Voice data used to train AI is predominantly from English-based languages. Languages other than English, regional accents, dialects, and non-dominant linguistic communities are absent or marginalised, leading to biases like those found in image datasets. AI cannot speak with diverse voices because it does not hear them. The LibriSpeech dataset – approximately 1000 h of 16 kHz read English speech from audiobooks predominantly from nineteenth century and earlier literature – was the most widely used voice-training dataset for two decades. A recent study found almost twice as many word error rates when AI listened to Black subjects speaking in African-American Vernacular English than for White speakers (Koenecke et al. 2020 p.7685). Counter datasets are being developed to address this audio representational narrowing. Mozilla's Common Voice project is a crowdsourced alternative; its goal is to provide a free, open-source database of recordings of voices embodied by speakers of varying ages, genders, and accents. As of June 2024, the dataset encompasses 31,841 h of recorded speech across 129 languages, with 20,789 h validated by the community. (Common Voice 2024 n.p.)

In one of their series of studies about voicing TTS, Ido Ramati presents TTS as 'Algorithmic Ventriloquism' (Ramati 2024, p.3), emphasising how AI cloned voices, trained on a specific person, are dissociated from that individual's embodied voice that trained them. When people experience dissociation, they disconnect from their thoughts, feelings, memories or sense of identity. The dissociation triggered by hearing a voice emanating from the 'wrong' body predates AI voice generation.

4 Methodology/analysis

4.1 Positionality statement

This paper is situated (Haraway 2013 p.589) and shaped by the author's embodied identity as a late middle-aged

cis-gendered heterosexual White British woman who lives with pain but is mainly well. This paper emerges from her experience of being located in and influenced by values of the Global North where she has lived for most of her life. She grew up in the middle of England and went to speech and drama classes at eight to reduce her regional accent. She is a feminist and first-generation graduate of fine art who, while not born digital, attended the UK's first masters programme in the late 1980s where artists and designers could use computers. She often works with scientists from a range of disciplines. Her professor role and the privilege of being an academic gave her access to resources to conduct this work. These intersectional perspectives and her shifting subject-position in relation to AI shaped the practice-based research and its interpretation.

Each method centres the researcher's identity and experience, and as such, they are best discussed in the first person. This paper has (e)merged with the production of *herbAirium*, a self-portrait video with a voice-over made in autumn 2024 in collaboration with generative AI. I trained the AI with photos of me, some of which were selfies, and three recordings of my voice. Through a combination of desk-based research and reflexive practice, I experienced, examined and countered the resulting lack of diversity, biases and classification challenges in AI instruments. The methodology for this study combines autoethnography, reflexive practice and autotheory. Autoethnography, which centers writing, lends itself to studying my production of AI-generated images. Firstly, because they are generated in response to my written text prompts. Secondly, the video's voice-over began as a written script, read by a synthetic voice trained on my speech. Thirdly, the overarching video narrative is an autoethnographic piece that mirrors my lived experience while creating an AI video as reflected in the first-person voice-over. Given the practice-based nature of my research, reflexive research methods helped to slow the making process down so that I could critically engage with both the process of AI generation and its implications. Since this reflexive engagement occurred through interactions with new instruments of vision and sound – generative AI – autotheory provided a valuable complementary framework. It prompted me to shape my theories about AI's technological, aesthetic, and epistemic constraints by engaging in AI-mediated self-representation.

4.2 Autoethnographic writing and speaking

The autoethnographic research method, which emerged from anthropology in the 1970s, involves a researcher writing about a personally relevant topic, drawing on tenets of autobiography. The researcher situates their experiences within their context, then systematically analyses those experiences to understand broader social, cultural, and political issues.

While the “graphy” signals the method's basis in written research, over the past 25 years visual artists have expanded autoethnography through practice-based approaches (Ellis et al. 2011 n.p.). The “auto” refers to the researcher's own experiences, making it valuable for artist-researchers by offering a framework to systematically analyse personal stories central to much artistic practice. Brydie-Leigh Bartleet describes performing autoethnographies as involving “the construction of accompanying twin narratives, the vulnerability of subjecting one's creative practice to scrutiny, and balancing artistic and aesthetic concerns with the rigors of research process” (Bartleet 2021, p.139). The “ethno” indicates the cultural context or group being examined, which can be further probed using Haraway's idea of “situatedness.” Autoethnographic researchers merge autobiography and ethnography to both make and write autoethnography. As a method, it is simultaneously process and product (Ellis et al. 2011, n.p.). It is particularly apt here as an embodied mode of inquiry, with this research centred on the artist's body. The AI-generated images and voice of that body are themselves autoethnographic products. The embodied nature of this research extends beyond written reflection on physical or emotional responses to those representations, as it unpacks embodiment in the creation of generative AI tools.

In this case, researcher writing is twofold: writing prompts and a voice-over script. This writing is essential to ‘making’ video content that explores personal and conceptual dimensions of digital twin relationships. The autoethnographic process is iterative– the researcher applies a method, observes its outcomes, reflects on those outcomes, and refines the approach accordingly. This learning curve manifests in the refinement of text prompts in AI image and video generation. After assessing the visual output of a generated image or video, the researcher revised the prompt to experiment with alternative inputs. RunwayML's system design inherently facilitates this reflexive cycle – it limits users to generating two 5- to 15 s video generations at a time. Frequent system slowdowns cause all users to wait several minutes before the resulting videos are available for evaluation. This enforced delay creates a built-in pause for reflection, reinforcing the iterative nature of prompt refinement and generative AI interaction.

Developing skills required for effective prompt engineering by crafting text inputs to guide AI-generated images and videos is challenging. One key difficulty is the lack of transparency in AI-generated content shared online, as these images and videos are rarely accompanied by the prompts that produced them, partly because when artists use publicly available generative AI instruments of vision (as opposed to those trained on custom datasets) written prompts are often seen as the creative and original act that results in images and therefore deliberately obfuscated. This is explained by AI researcher Ethan Smith, who shares prompts (and whose

prompt guides I referred to) as part of their commitment to “learning in public” (Smith 2024 n.p.). Smith notes “the only things that separate our work from that of others is our prompts, and our settings. [...] For that reason, I understand why some AI Artists in the community aren’t open to sharing information.” (Smith And Lam 2022 n.p.) Despite this, some prompts connected to images produced by artists are well-documented. I referred to guides and templates for prompting developed by groups working together to solve problems shared in online Discord groups and found in the extensive resource compiled by Smith, *A Traveler’s Guide to the Latent Space* (Smith And Lam 2022 n.p.). Prompt engineering has been described as resembling a dialogue with text-to-image generative AI systems, yet despite its centrality to AI image and video generation, it is an emerging and under-documented practice. Oppenlaender, whose prompting guides I also accessed, notes that AI prompting “still resembles a cottage industry, with concepts and structures yet to emerge.” (Oppenlaender 2024a p. 3771) The lack of standardised frameworks makes prompt engineering an experimental, iterative process, reinforcing the need for hands-on engagement and reflective adaptation in practice-based research. This is especially true for writing texts that are combined with a starter image to create video clips.

4.3 Autotheory

My method of producing *herbAlrium*, including prompt writing, is an example of what has recently been termed “autotheory”, “taking one’s embodied experiences as a primary text or raw material through which to theorize, process, and reiterate theory to feminist effects.” (Fournier 2018 p.646). Autotheory is built upon precursor conceptual, performance, and body art practices, especially those used and theorised by Chicana feminists like Gloria Anzaldúa, whose creative writing articulated intersectional realities and proposed new theories through works such as *Borderlands* (Woodward 1989 p.531). Black feminist epistemology emphasises the power of using dialogue to assess knowledge claims and interacting with the object of knowledge rather than observing it from a detached distance, (Collins 1990 p.550), exemplified in Audre Lorde’s nuanced descriptions of her cancer treatment. Lorde includes verbatim recollections of comments that situate her treatment as racist and misogynistic in *The Cancer Journals*, (Lorde 1997 p.59). Extending Lorde and bell hooks’ work, Ellen Samuels emphasises that autotheory and its associated art practices are too often erroneously centred on White Well Women. In *Twenty-Seven Ways of Looking at Crip Autotheory* (Samuels 2023 p.203), she counters this with the assertion that “There is no theory of autotheory that does not start with the ill and disabled bodymind. There is no theory of autotheory that is

not already crip.” (Samuels 2023 p.203), reminding us that autotheory is a theory and practice of embodiment.

HerbAlrium is at the heart of this paper, comprising AI-generated video and audio clips that represent moments from my life, featuring my digital twin’s face and body intercut with other AI-generated video clips that depict my art studio and the nature walks integral to my artmaking practice. Combined with an AI voice-over, this is a ‘meta’ piece that critically materialises experiences of using AI to generate images of an identical-looking digital twin called “Jane”. This identical-looking digital twin spends most of her time working as an artist and amateur botanist (as opposed to my real life, in which I often present as an academic or carer). The video’s final version includes AI-generated clips that failed to present an externalised form that I could recognise as looking like me, with others that looked more like me. The AI voice-over, spoken by my cloned voice, reads a script that includes reflections on creating the video. I use autotheory in the discussion section in a straightforward account of how I learned about the infrastructure of AI image generation by utilising it to create this new video artwork.

Lauren Fournier connects autotheory and performance, highlighting the expediency of performance as a medium, with their bodies a material that artists have easy access to. By contrast, the video artwork described here was made using less accessible material, generative AI paid for by a subscription costing \$28 per month, accessible due to my privilege as a salaried academic. In autotheory, personal embodied experiences serve as the foundation for theoretical exploration and the lens through which those experiences are interpreted and articulated. Fournier’s book, *Autotheory as feminist practice in art, writing, and criticism*, introduced me to the work of Dutch cultural theorist, curator, and video artist Mieke Bal who explores autotheory in “*Documenting What? Autotheory and Migratory Aesthetics*” (Bal 2015 p.124), defining it as both a practice and an ongoing, spiralling dialectic between analysis and theory. For Bal, creating documentary films and then analysing them through a theoretical lens bridges artistic practice and critical inquiry. My approach owes much to Bal’s dialogic shifts between theoretical analysis and creating. However, my approach differed from Bal’s as my dialogic shifts took place throughout the making process, clip by clip, reflexively, rather than after the time-based work was complete. Autotheory centres the embodied practitioner in the process of making self-images. The positionality statement situates the research by revealing more about the author’s identity. This is not to suggest that we can, or would want to, rise above the baggage and attempt to become ‘objective’, but that becoming aware of and accounting for our preconceptions, how our identity and lived experiences influence us, can enrich our planning, conducting, evaluating, and sharing our research.

4.4 Practicing reflexively

This paper and the artwork have emerged through a “reflexive practice” (Argyris And Schon 1974). The reflexive practitioner becomes usefully ‘self-conscious’ and makes themselves aware of the baggage and experiences they bring to research. Making *herbAIrium* was a real-time practice. To learn how to use AI technological infrastructure, I drew on Chris Argyris and Donald Schön’s ideas for how professional practitioners succeed in learning by “developing one’s own continuing theory of practice under real-time conditions” (Argyris And Schon 1974, p.157). While producing *herbAIrium*, I became increasingly reflexive ‘in the moment’ of research, evaluating AI images generated in response to prompts and my emotional, psychological and physical reactions to them. I began to develop theories about the representation of gender, race, and age in AI-generated images. The process of reflection-in-action is essentially artistic; the practitioner makes judgments and exercises skills for which no explicit rationale has been articulated, but in which she nevertheless feels an intuitive sense of confidence, (Brookfield 1986 p.247). I was mindful of surprise, puzzlement, or confusion during this real-time research because such reactions often signal something of value to our theories that we might dismiss as ‘noise’ in our data or process. For example, in struggling to create a text prompt to produce an image that corresponded to a storyboarded idea, I first dismissed some images as “failures” but, remembering to be reflexive in the moment, I looked at them to see what they could tell me about the prompt I had written, what could be inferred about the datasets the ML had been trained on.

4.5 Cloning my voice for my digital twin

Like most speaking people, my voice changes depending on how tired I am and who I am conversing with. My accent is relational and changes as I code-switch, converging to align with voices around me or diverging when I want to distance myself. It has also changed significantly throughout my life. I was born in Birmingham in the UK and I developed a Birmingham, or “Brummie” accent, which linguistic studies have consistently shown “engenders the most negative connotations because it is widely viewed as an ‘incorrect’, ‘ugly’, ‘common’ and ‘uneducated’” (Thorne 2005 p.154). If one of feminism’s goals is to have a voice that carries authority so that (literal) speech acts have the outcome we intend, then that speech act needs to be recognised. What happens when, as a speaker, we experience accent discrimination and our voices cause people to diverge from us, and our speech acts are not recognised?

I was eight when my family moved, and my new classmates claimed not to understand me and ridiculed my accent. My speech did not have its intended effect, and I

was unable to integrate. My mother enrolled me in speech and drama classes to help me develop a different accent and a ‘better’ posture via repeated vocal and performative practices (Burchell 2024 p.356). This was the 1970s, when Britain standardised elocution and drama training stressed clarity and focused on reinforcing the musculature of the organs involved in speech (Burchell 2024 p.360) to change the body to “improve” the voice. My accent (and posture) changed. When I listened to my voice clone for my digital twin, I cringed—a response shared by many people when they hear recordings of their own voice, but not necessarily because of any perceived accent. This is because we never hear our own voices in the way others do when we talk. Our voice recordings usually sound higher-pitched than we expect when we listen to ourselves speak (own-voice). Our reaction to a recording of our voice typically sounds both familiar and eerie (Kimura And Yotsumoto 2018 p.12).

Though hearing a recording of our voice might make us cringe, it is nothing compared to the alienation of listening to an off-the-shelf AI synthetic voice speak ‘for’ us. At the time of writing, out of RunwayML’s 19 pre-trained voices, 11 are categorised as “feminine”, one of which is British and a second “English-Swedish”, and neither sounded enough like me for me to want to use. The two voices labelled as “pleasant” are both tagged “feminine”, as are individual voices described as “seductive” and “gentle”. This collection’s only “authoritative” voice is a “masculine” voice. Like many generative AI voices, those provided by Runway ML reproduce gender, ethnicity and class stereotypes and culturally code them into pre-defined voice characters. It is not surprising that I chose to clone my voice to make what RunwayML describes as a “custom voice” to use with my digital twin and for voice-overs.

5 Discussion

5.1 Situating fieldwork with RunwayML

Working reflexively with RunwayML involved systematic observations, which were recorded in handwritten notes. Some were adapted for inclusion in the voice-over script. During my iterative learning cycle, I refined my prompt engineering. Because RunwayML, like most generative AI instruments, is changing very rapidly I added a section to the video voice-over to identify when the video was made, “I made this piece in 3 weeks, during August and September 2024 [...] Being flexible also allowed me to try new techniques as they came on stream, like the option to extend one 5 to 10 s sequence, which became available in early September 2024.” (unpublished voice-over script).

RunwayML’s environment was developed with Stable Diffusion, a deep learning generative artificial neural

network, designed to generate images based on a user's text descriptions and accessed via a commercial API (Birhane et al. 2023 n.p.). Stable Diffusion was trained on LAION-5B's publicly available dataset derived from data scraped from the web. Almost half of LAION's image-alt-text pairs came from only 100 domains. Pinterest comprised 8.5% of the subset, followed by websites such as WordPress, Blogspot, Flickr, DeviantArt and Wikimedia where many image-alt-text pairs, as described earlier, are incongruent. AI models trained on flawed pairings may struggle with accurate semantic alignment, fail to connect visual and textual elements meaningfully, or generate offensive, misleading, or harmful images when given benign text prompts. These insights might help explain some of the challenges I encountered when using text and image prompting in RunwayML.

I used prompt guides from outside Runway, as discussed above, and engaged with Runway Discord prompting groups. As of January 2025, Similarweb data shows about 38% of visitors to Runway's website are "female" and 62% are "male". However, these figures may not reflect the actual gender makeup of Runway's users, as the company has not released official user data.

Central to my generative AI video workflow was creating a digital twin named Jane. Initial text-to-image prompts in RunwayML failed to generate a character resembling me closely enough to qualify as a self-portrait. I therefore trained a custom character using photographs of myself. RunwayML permitted 15–30 still images for training, recommending cropped head-and-shoulder shots in varied lighting. As a result, the character was based on a limited set of images showing only part of my body—the head, which typically makes up one-eighth of the human form. After uploading the photos, you name the character and include that name in your text prompts, while selecting it from a dropdown menu. I named mine "Jane" to reflect its intended role as my identical-looking digital twin. Then, in RunwayAI's Gen-3 Alpha, I started to explore prompts such as "Medium close-up looking at the camera, scrubland in the background. Jane."

While much AI prompting literature focuses on text-to-image generation, I had to develop skills that combined text prompts with a starting image for video expansion. Guidelines for this type of prompt engineering are limited. RunwayML provides some documentation, which was helpful but insufficient. To compensate, I engaged daily with RunwayML's Discord community—reviewing clips, joining discussions, and testing shared techniques. Although users are encouraged to post both prompts and outputs for collaborative troubleshooting, most withheld their text inputs. This limited peer learning, though it's understandable—many artists regard their prompts as core to their creative practice, sometimes even more than the images they produce.

5.2 Name discrimination

In retrospect, using the label "Jane" when I saved my custom character may have reinforced AI bias through name discrimination. Research by social scientists has shown race- and gender-based name discrimination in the labour market (Bertrand And Mullainathan 2004 p.10), with candidates with names associated as being White and masculine being privileged over those with ethnic-sounding names when CVs are filtered using AI. Large Language Models (LLMs) have also perpetuated humans' discrimination based on perceived race or ethnicity associated with names (An et al. 2024 p.391). An analysis of bias related to given names shows that GPT-3.5-Turbo and Llama 3 prefer candidates with White-aligned given names and suggest higher salaries for those candidates (Nghiem et al. 2024, p.7272). So how might LLMs categorise the name "Jane"? An online search on associations with the given name "Jane" reveal it is associated with White women over the age of 50 (teacoffeecream 2023 n.p.), categorise it as "posh" (Green 2025 n.p.), and the Name/Nerds group of Reddit users describe it as "Slightly old fashioned, British-sounding" (mattymillyautumn 2017 n.p.), which resonate with many of the elements in my earlier positionality statement but diverges in others.

5.3 Gender representation and bias in AI-generated images

Nettrice Gaskins states "The secret lies in the prompt, more specifically in what words are used in the prompt." (Gaskins 2023a n.p.) As a cisgender heterosexual woman who does not always present as conventionally feminine, I encountered persistent gendered biases in AI-generated images. When generating my digital twin, Jane, attempts to create a neutral, androgynous-looking character were hindered by what appeared to be the AI model's binary gender defaults. Prompts without gender descriptors produced masculine-presenting figures. Using the iterative, experimental approach outlined in prompt engineering studies (Oppenlaender 2024b p.333), I adjusted my prompts by adding gendered terms. Words like feminine produced hyper-feminised results, and woman yielded images with exaggerated, conventionally sexualised traits—such as an overstated hourglass figure—that did not align with my body image (see Fig. 1).

To counter AI's bias toward generating White-presenting figures, researchers created instructions for GPT-3 to produce more inclusive prompts. These directed it to describe people in specific occupations with diversity, detail, imagination, and emotion—always including the person's ethnicity (Clemmer et al. 2024, p. 8599). While writing prompts, I applied this approach to integrate gender information. Terms like androgynous had little effect, but androgenic

Fig. 1 Four different images of the author's digital twin that show a range of gender presentations using “androgenic” in the text prompt. AI-generated images created by RunwayML in September 2024 via a text prompt applied to the author's custom character that was trained with photos of the author



produced more accurate results—feminine-presenting bodies with narrower hips and smaller breasts that aligned more closely with my intended digital twin. Jane's body varied across outputs, though never appeared with a visible disability, and my Whiteness was consistently rendered. Since Jane was trained only on head-and-shoulders images, the absence of full-body references may explain the inconsistency in body shape. Yet, privileging the face should mean facial features better reflect the training images—this wasn't the case. AI's limitations in rendering gender-diverse bodies became clearer when fine-tuning specific traits. For instance, despite being trained on images showing my brown eyes, the model often generated blue or green ones. Even when “brown eyes” was specified in prompts, lighter eye colours persisted. This points to dataset biases, particularly linked to the tag “woman”, which favour stereotypical colonial and Western beauty norms, including Whiteness and idealised features like blue eyes (Abdul Kader 2020, n.p.).

5.4 Struggles with hands: feminized defaults in AI training data

My digital twin's AI-generated hands posed a persistent challenge. Unlike the familiar issue of extra fingers or distorted anatomy (Matthias 2025, n.p.), my frustration came from repeated depictions of long, manicured, painted nails—unlike my short, unpainted ones. Despite extensive prompting, the AI defaulted to hyper-feminised aesthetics. *Short bare dirty fingernails* yielded long bright red nails; a second attempt produced matte black ones. *Unpainted nails* resulted in frosted pink points, and *gardener's hands* returned white painted tapered nails. At least one finger always wore a ring—usually on the traditional marriage finger—which I

couldn't control; prompting often led to more rings. Even *dirty fingernails* prompted a video of manicured nails with dirt at the cuticles. Workarounds like *wearing worn-out gardening gloves*, meant to conceal the nails, still showed long, frosted pink nails before finally generating a usable image of my digital twin in gloves. I addressed this visually and verbally in a voice-over: “Hang on! That's not right, I don't wear nail varnish, and why has the AI given me a wedding ring automatically? Let me change the prompt to dirty fingernails. Now the nails are black. Unmarried? Wow, even longer nails and more rings!” (unpublished voice-over script) (see Fig. 2).

Discussions in RunwayML's Discord suggest this persistent bias stems from training datasets likely scraped from platforms like Pinterest and Instagram. This could explain why my AI twin's hands consistently reflected hyper-feminised beauty standards—not reality, but a curated, commercialised vision of femininity embedded in the data. To integrate these autotheory insights into the final video, I added a voice-over alongside clips visualising the underlying database: “*I guess to create these images, the Runway AI analyzes the starter image I provide and my text prompts, extracts features and keywords, and compares them to templates in its database, its taxonomy of images and texts, to generate new images. Those databases are limited and perpetuate biases. Women's hands have wedding rings, and even dirty fingernails are hidden beneath a perfect manicure.*” (unpublished voice-over script).

Similar biases appeared in clothing and props. When prompting Jane to carry a practical bag—*rucksack*, *backpack*, or *worn holdall*—the AI rendered her in a boiler suit and Wellingtons, holding designer handbags like trophies. Attempts with prompts like *torn plastic*, *paper carrier bag*,

Fig. 2 Four examples of AI-generated hands, each showing a character with painted fingernails. Made with RunwayML



or *cloth bag* also failed. The intended aesthetic was never realised. Clothing defaults consistently included tight jeans and fitted tops, reinforcing feminine-coded norms. To counter this, I used prompts like *scruffy boilersuit with frayed cuffs, worn-out gardening gloves, and frayed fabric*, which more successfully de-emphasised gendered fashion tropes (see Fig. 3).

5.5 Identity and dissociation in digital twin representation

Perhaps the most alienating aspect of creating AI-generated video was the lack of consistency in Jane's appearance across sequences. No matter how often I fine-tuned my prompts, Jane never consistently resembled me. Even when using the same starting image—a still from the Jane custom character chosen for its likeness—and the same text prompt, each iteration produced a different-looking figure.

Confronted with this inconsistency, I changed tack. Instead of generating scenes that matched my storyboarded

Fig. 3 Two examples of AI-generated images of the author's custom character holding a bag. Made with RunwayML



shots—long shots, close-ups, big close-ups—I adopted a new strategy: selecting a single clip where Jane appeared on camera, then generating the rest from her point of view, where she would not be visible. I spent hours watching clips of various Janes, trying to choose one. At my standing desk with a large monitor, Jane, in close-up, appeared life-sized, meeting my gaze. I found myself staring into her shifting, AI-generated eyes—brown, green, most often blue—searching for something familiar.

Then something unsettling happened. After hours looking at my AI twin, when I stepped away for a cup of tea or food, I would catch my reflection in the mirror near the studio door—and not immediately recognise myself. There was a fleeting disconnect from my own face. I had spent so much time studying, adjusting, and identifying with these digital twins that my reflection felt unfamiliar for a moment. I began noticing subtle features about my face and body—small details I had never observed before.

5.6 “Default smile” and gendered expectations

One of the things I noticed was my neutral facial expression. Watching AI clips of my digital twin, I felt a slight upward twitch at the corners of my mouth. She was almost always shown smiling, triggering my unconscious tendency to mirror facial expressions. Noting this, I began occasionally checking my reflection in the computer screen to see what I looked like when concentrating. Consistent with research showing that “women” generally smile more than “men,” but that these gender differences disappear after late middle age, aged 60, I wasn’t smiling when focused. The AI’s persistent addition of a slight smile to Jane’s face raised further questions about gender norms in visual datasets. Even when prompted with terms like *serious*, *frowning*, or *distracted*, Jane’s expression often returned to what I interpreted as a subtle, flirtatious smirk. This aligns with research on gender and emotional expression in AI training datasets, where feminine-presenting faces are more frequently depicted smiling than masculine-presenting ones. I speculated about whether this default gender presentation was influenced by the dataset’s likely composition, which skews toward younger, socially curated images – often drawn from selfies, fashion, and social media platforms where most humans smile for the camera. My experience mirrors broader research on gendered ageing and social perceptions of emotion that suggests that older individuals’ neutral facial expressions are often misinterpreted as sadness or grumpiness, especially by younger observers. If the AI’s training data predominantly features younger, smiling faces, this may explain why it struggles to represent older or neutral expressions without reinforcement. In this sense, AI’s “default smile” not only perpetuates gender biases in representation but also aligns

with age-related aesthetic norms that favour youthfulness and upbeat emotional performativity.

5.7 Voice

Although speaker identification depends on linking voice to person, I felt less dissociation than expected when hearing my cloned voice from the ‘wrong’ body during lip-syncing a masculine version of my digital twin. This might be because, following feminist pedagogy, ‘voice’ is a metaphor for (textual) authority (Haydari et al. 2023). Using RunwayML’s slightly clunky lip sync feature, my digital twin found her/their voice; while presenting differently she/they spoke with one voice. The synced voice made me identify more with versions I had felt alienated from as mere moving images. Consequently, I re-edited the video to begin with multiple gender presentations of my digital twin, accompanied by a voice-over in my cloned voice: “Sometimes, it’s hard to recognise myself in the images that AI generates based on the character template trained with photos of me. Still, I hope that you’ll recognise my AI voice throughout this video despite the appearance of my character changing.” (unpublished voice-over script).

6 Conclusion

AI actively participates in the materialisation of bodies – who is made visible, who is excluded, and how gender, race, disability, and other identity markers are algorithmically shaped. This aligns with Barad’s critique of representationalism, extended here to consider generative AI as an active participant in the ongoing process of materialising bodies, rather than as a neutral tool that reflects reality.

Interrogating the interplay between technology, the body, and the processes of digital self-representation that reinforce or disrupt socially constructed notions of gender, ageing, and race reveals how the artist’s self-portrait or digital twin prompts feelings of alienation, recognition, and empathy. AI models trained on flawed pairings may struggle with accurate semantic alignment, fail to connect visual and textual elements meaningfully, or generate offensive, misleading, or harmful images when given benign text prompts. The practice described here, of making an AI-generated video featuring an identical-looking digital twin, revealed some of those deep-seated biases in AI-generated presentations, particularly regarding gender, race, and ageing. The constraints I encountered—such as hyper-feminised body shapes, persistent gender defaults, and the difficulties in generating neutral or serious facial expressions—highlight implicit norms embedded in AI training data. These biases not only shape how digital identities are constructed but also

reinforce dominant cultural aesthetics that prioritise youth, hyper-femininity, and rigid binary gender presentations.

However, the paper also shows that artists can, to an extent, counter these biases by curating and using our own AI characters and voices that are more representative. We can create our own datasets to train AI and utilise existing AI instruments for vision to generate intentionally imperfect, diverse, or exaggerated representations. Both these strategies draw attention to areas where AI models fail or misinterpret data and expose systemic issues within the collaborative process. The broader implications of these findings extend beyond individual artistic experimentation. They emphasise the need for critical interrogation of AI training datasets, greater transparency in dataset composition, and the development of more inclusive AI tools that account for the full spectrum of human diversity and self-representation.

More work is needed to counter AI bias, especially by artists with intersectional identities, which necessitates removing barriers to accessing AI. Dr. Nettrice R. Gaskins drew attention to the urgent need for more tools, films, and art to “channel alternative frameworks that address equity through the generation of new ideas and prototypes that counter bias and other negative effects of AI in underrepresented and under-resourced groups,” (Gaskins 2023b, p. 423), and that project is of ongoing urgency.

Author contribution The paper was sole authored.

Data availability No datasets were generated or analysed during the current study.

Declarations

Conflict of interest The authors declare no competing interests.

Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

References

- AbdulKader S (2020). Beauty is in the eye of the colonizer: eurocentrism rears its “beautiful” head. In: Medium. <https://medium.com/@salmaabdulkader/beauty-is-in-the-eye-of-the-colonizer-eurocentrism-rears-its-beautiful-head-519b8acff6e4>. Accessed 14 Mar 2025
- An H, Acquaye C, Wang C, et al (2024). Do large language models discriminate in hiring decisions on the basis of race, ethnicity, and gender? ArXiv Prepr ArXiv240610486
- Argyris C, Schon D (1974). Theory in practice: inquiry into a profession of learning
- Arnold K, Bandelli A, Ellis R, et al (2024). Ten out of ten: a review of the last decade. Sci Mus Group J Autumn 2024: <https://journal.sciencemuseum.ac.uk/article/ten-out-of-ten-a-review-of-the-last-decade/#text-5>. Accessed 15 Mar 2025
- Atairu M (2023). Portraits of mami wata: blonde braids. In: Minne Atairu. <https://minneatairu.com/blondebraids>. Accessed 15 Mar 2025
- Bal M (2015). Documenting what? Auto-theory and migratory aesthetics. Companion Contemp Doc Film 124–144
- Barad K (2007) Meeting the universe halfway quantum physics and the entanglement of matter and meaning. Duke University Press, Durham
- Bartleet BL (2021) Artistic autoethnography: exploring the interface between autoethnography and artistic research. Handbook of autoethnography. Routledge, pp 133–145
- Bertrand M, Mullainathan S (2004) Are Emily and Greg more employable than Lakisha and Jamal? a field experiment on labor market discrimination. Am Econ Rev 94:991–1013
- Birhane A, Han S, Boddeti V, Luccioni S (2023) Into the laion's den: investigating hate in multimodal datasets. Adv Neural Inf Process Syst 36:21268–21284
- Birhane A, Prabhu VU, Kahembwe E (2021). Multimodal datasets: misogyny, pornography, and malignant stereotypes. ArXiv Prepr ArXiv211001963
- Brookfield S (1986) Understanding and facilitating adult learning: a comprehensive analysis of principles and effective practices. McGraw-Hill Education (UK)
- Buolamwini J, Gebru T (2018) Gender shades: intersectional accuracy disparities in commercial gender classification. PMLR, pp 77–91
- Burchell A (2024) The art of speech: elocution, speech training, speech therapy, and the performative limits of class in mid-twentieth-century Britain. Mod Br Hist 35:354–374. <https://doi.org/10.1093/tcbh/hwae043>
- Clemmer C, Ding J, Feng Y (2024). Precisedebias: An automatic prompt engineering approach for generative AI to mitigate image demographic biases. pp 8596–8605
- Collins PH (1990) Black feminist thought in the matrix of domination. Black Fem Thought Knowl Conscious Polit Empower 138:221–238
- 'Common Voice 18 Dataset Release' (2024). Available online: <https://foundation.mozilla.org/en/blog/common-voice-18-dataset-release/> (Accessed 13 March June 19, 2024).
- Crawford K, Paglan T. Excavating AI: the politics of images in machine learning training sets. In: Excav. AI. <https://paglen.studio/2020/04/29/imagenet-roulette/>. Accessed 15 Mar 2025
- Crenshaw K (1989). 'Demarginalizing the intersection of race and sex: a black feminist critique of antidiscrimination doctrine, feminist theory and antiracist politics', University of Chicago Legal Forum, 1989(1): 31.
- d'Inverno M, Prophet J (2006) Biology, computer science and bioinformatics: Multidisciplinary models, metaphors and tools, multidisciplinary investigation into adult stem cell behaviour. Transactions on computational system biology. Springer, London, pp 49–64
- Dinkins S. Not the Only One. In: Stephanie Dinkins. <https://www.stephaniedinkins.com/ntoo.html>. Accessed 15 Mar 2025
- Eifer M (2020). Prosthetic Memory. In: BlinkPopShift. <http://www.blinkpopshift.com/project-pages/prosthetic-memory>. Accessed 15 Mar 2025
- Ellis C, Adams TE, Bochner AP (2011) Autoethnography: an overview. Hist Soc Res/Historische Sozialforschung 36:273–290

- Fournier L (2018) Sick women, sad girls, and selfie theory: autotheory as contemporary feminist practice. *Auto/Biogr Stud* 33(3):643–662. <https://doi.org/10.1080/08989575.2018.1499495>
- Gaskins N (2023b) Interrogating algorithmic bias: from speculative fiction to liberatory design. *TechTrends* 67:417–425. <https://doi.org/10.1007/s11528-022-00783-0>
- Gaskins N (2023a). The role of the artist in generative AI. In: Medium. <https://nettricegaskins.medium.com/the-role-of-the-artist-in-generative-ai-de79d6b8098e>. Accessed 3 Mar 2025
- Green C (2025). The 100+ poshest names in Britain. In: Nameberry. <https://nameberry.com/blog/the-100-poshest-names-in-britain>. Accessed 3 Mar 2025
- Hancox-Li L, Kumar IE (2021). Epistemic values in feature importance methods: Lessons from feminist epistemology. pp 817–826
- Haraway D (2013) Situated knowledges: the science question in feminism and the privilege of partial perspective 1. *Women, science, and technology*. Routledge, pp 455–472
- Haydari N, Sesigür O, Ulutaş A, Irmak B (2023) Sound as technologies of the self for feminist pedagogy. *Gend Educ* 35:421–436
- Hayles NK (2004) Print is flat, code is deep: the importance of media-specific analysis. *Poet Today* 25:67–90
- Hemment D, Currie M, Bennett SJ, et al (2023). AI in the public eye: investigating public AI literacy through AI art. pp 931–942
- Prophet J (2024) herbAIrium, [AI generated video] available online:<https://vimeo.com/1082571147> Password = 5984crow.
- Jones D, Snider C, Nassehi A, Yon J, Hicks B (2020) Characterising the digital twin: a systematic literature review. *CIRP J Manuf Sci Technol* 29:36–52
- Kimura M, Yotsumoto Y (2018) Auditory traits of “own voice.” *PLoS ONE* 13:e0199443
- Koenecke A, Nam A, Lake E et al (2020) Racial disparities in automated speech recognition. *Proc Natl Acad Sci* 117:7684–7689. <https://doi.org/10.1073/pnas.1915768117>
- Liao C (2024) My avatar’s avatars: a visual exploration and response to AI-generated avatars. *Vis Cult Gend* 19:11–23
- Lorde A (1997). The cancer journals [1980]. *J Lond Sheba Fem* 1985
- Matthias M (2025). Why does AI art screw up hands and fingers? | Explanation, tools, & facts | Britannica. <https://www.britannica.com/topic/Why-does-AI-art-screw-up-hands-and-fingers-2230501>. Accessed 14 Mar 2025
- Mattymillyautumn (2017). Slightly old fashioned, British-sounding names? In: r/namenerds. www.reddit.com/r/namenerds/comments/7cs0j8/slightly_old_fashioned_britishsounding_names/. Accessed 4 Mar 2025
- Mills M (2012) The audiovisual telephone: a brief history. *Handheld? music video aesthetics for portable devices*. ART-Dok, pp 34–47
- Nghiem H, Prindle J, Zhao J, Daumé III H (2024). “You Gotta be a Doctor, Lin”: an investigation of name-based bias of large language models in employment recommendations. *ArXiv Prepr ArXiv240612232*
- Oppenlaender J (2024a) A taxonomy of prompt modifiers for text-to-image generation. *Behav Inf Technol* 43:3763–3776
- Oppenlaender J (2024b) The cultivated practices of text-to-image generation. *Humane autonomous technology*. Palgrave Macmillan, Cham, pp 325–349
- Prophet J (2001) *TechnoSphere*: “real” time, “artificial” life. *Leonardo* 34:309–312
- Prophet J, Kow YM, Hurry M (2018) Small trees, big data: augmented reality model of air quality data via the Chinese art of ‘artificial’ tray planting. *ACM SIGGRAPH 2018 posters*. ACM, New York, NY, USA
- Prophet, J (2024) herbAIrium. <https://vimeo.com/1082571147/de22e3e781?share=copy> Accessed 14 Mar 2025
- Ramati I (2024) Algorithmic ventriloquism: the contested state of voice in AI speech generators. *Soc Med + Soc* 10:20563051231224400
- Redler-Hawes H. Hannah Redler-Hawes. In: Open Data Inst. <https://theodi.org/profile/hannah-redler-hawes/>. Accessed 15 Mar 2025
- Ridler Mosaic Virus, 2019. In: Anna Ridler. <http://annaridler.com/mosaic-virus>. Accessed 15 Mar 2025
- Samuels E (2023) Twenty-seven ways of looking at crip autotheory. *Crip authorship*. New York University Press, pp 203–209
- Scheuerman MK, Wade K, Lustig C, Brubaker JR (2020) How we’ve taught algorithms to see identity: constructing race and gender in image databases for facial analysis. *Proc ACM Hum-Comput Interact* 4:1–35
- Smith E, Lam J (2022). A traveler’s guide to the latent space. <https://www.notion.so>. Accessed 14 Mar 2025
- Smith E (2024). Research blog. In: Ethansmith2000. <https://www.ethansmith2000.com/blog-posts>. Accessed 14 Mar 2025
- Sobande F (2021) Spectacularized and branded digital (re) presentations of black people and blackness. *Telev New Media* 22(2):131–146
- Sterne J (2012) *The sound studies reader*. Routledge
- Sun L, Wei M, Sun Y et al (2024) Smiling women pitching down: auditing representational and presentational gender biases in image-generative AI. *J Comput-Mediat Commun* 29:zmad045
- Teacoffeecream (2023). Is the name Jane typically associated with white women? In: r/namenerds. www.reddit.com/r/namenerds/comments/18f3v0d/is_the_name_jane_typically_associated_with_white/. Accessed 3 Mar 2025
- Thorne S (2005) Accent pride and prejudice: are speakers of stigmatized variants really less loyal? *J Quant Linguist* 12:151–166
- Woodward C (1989). *Borderlands/La Frontera: The New Mestiza*, 530–532.
- Young E, Wajcman J, Sprejer L (2021). *Where are the women? mapping the gender job gap in AI*. London: The Alan Turing Institute. Alan Turing Institute Available online: <https://www.turing.ac.uk/news/where-are-women-mapping-gender-job-gap-ai>. Accessed 3 Mar 2025
1. Fredrickson, B. L. & Roberts, T. Objectification theory: Toward understanding women’s lived experiences and mental health risks. *Psychology of women quarterly* 21, 173–206 (1997).
 1. Mulvey, L. *Visual and Other Pleasures*. (Springer, 1989).
 1. Calogero, R. M. A test of objectification theory: The effect of the male gaze on appearance concerns in college women. *Psychology of women quarterly* 28, 16–21 (2004).
 1. Calogero, R. M. A test of objectification theory: The effect of the male gaze on appearance concerns in college women. *Psychology of women quarterly* 28, 16–21 (2004).

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.