



UWL REPOSITORY

repository.uwl.ac.uk

The algorithmic gaze: AI, attention and the ethics of brand perception

Olsen, Dennis (2026) The algorithmic gaze: AI, attention and the ethics of brand perception. Journal of Digital Media & Interaction, 9 (21). pp. 23-40. ISSN 2184-3120

<https://doi.org/10.34624/jdmi.v9i21.41168>

This is the Published Version of the final output.

UWL repository link: <https://repository.uwl.ac.uk/id/eprint/14769/>

Alternative formats: If you require this document in an alternative format, please contact: open.research@uwl.ac.uk

Copyright: Creative Commons: Attribution-Noncommercial-No Derivative Works 4.0

Copyright and moral rights for the publications made accessible in the public portal are retained by the authors and/or other copyright owners and it is a condition of accessing publications that users recognise and abide by the legal requirements associated with these rights.

Take down policy: If you believe that this document breaches copyright, please contact us at open.research@uwl.ac.uk providing details, and we will remove access to the work immediately and investigate your claim.

Rights Retention Statement:

The Algorithmic Gaze: AI, Attention and the Ethics of Brand Perception

Dennis Olsen
University of West London, United Kingdom
dennis.olsen@uwl.ac.uk
0000-0003-1911-7009

Received: 01-10-2025

Accepted: 04-03-2026

Abstract

This article critically examines how AI-supported visual evaluation tools – particularly predictive attention modelling software – are reshaping brand perception and design decision-making in hyperdigital environments. Drawing on a two-phase qualitative study involving practitioners and subject expert, it introduces the concept of the algorithmic gaze to interrogate how these systems simulate, encode and prescribe normative visual hierarchies. The analysis reveals that while such tools offer operational efficiency, they often obscure cultural nuance, misrepresent audience engagement and displace human interpretive agency. Practitioner interviews highlight tensions between usability and trust, especially among junior professionals, while expert reflections situate these dynamics within broader discourses of epistemic authority, ethical design and algorithmic bias. These insights are contextualised by the rhetorical framing of AI tools as scientifically authoritative, which shapes practitioner expectations and contributes to a culture of deference to data. The article synthesises its findings through the Interpretive Stewardship heuristic – a model comprising three interlinked dimensions that offer a scaffold for ethical engagement with AI tools in branding. It argues that branding must be re-framed as a site of cultural negotiation rather than computational optimisation and calls for regulatory, institutional and design-level commitments to transparency, inclusivity and reflexive practice. These findings contribute to emerging debates in branding, media theory and AI ethics.

Keywords *Algorithmic gaze, AI ethics, Infosphere, Interpretive Stewardship, Feminist data ethics, Brand perception*

1. Introduction

Artificial intelligence (AI) is rapidly transforming the landscape of brand communications, not only by automating design processes but by reshaping how visual meaning is constructed, evaluated and legitimised. In the UK, this transformation is unfolding within a broader national ambition to lead the global AI revolution. The recent UK Government's *AI Opportunities Action Plan* calls for a "push hard on cross-economy AI adoption", urging both public and private sectors to rapidly pilot and scale AI products to "drive better experiences and outcomes for citizens and boost productivity" (DSIT, 2025). Branding and advertising are central to this agenda, with AI-supported tools increasingly positioned as solutions for streamlining creative workflows and enhancing message effectiveness.

Among these tools, AI-supported visual evaluation solutions – such as predictive eye-tracking software – claim to simulate human attention with high fidelity, offering heatmaps, gaze plots and at-

tention scores that purport to guide creatives toward more effective brand messaging. Yet, as recent research demonstrates in a comparative study of research-grade eye-tracking and AI-supported software, such claims often mask significant discrepancies (Olsen, 2025).

Despite their growing adoption, there remains a critical gap in understanding how these tools influence interpretive agency, cultural representation and ethical decision-making in branding practice (Erickson, 2024; Heigl, 2025; Hue & Hung, 2025). Existing discussions tend to focus on technical performance or user convenience, often overlooking the normative assumptions embedded in algorithmic outputs and the epistemic consequences for practitioners. Moreover, while algorithmic bias and opacity have been widely examined in other domains, their implications for visual culture and strategic brand communication remain underexplored. This article addresses this gap by interrogating how predictive attention modelling tools shape brand perception and design decision-making in hyperdigital environments. It introduces the concept of the *algorithmic gaze* to explore how these systems encode and reproduce assumptions about attention, desirability and meaning. Drawing on empirical interview data and a discourse analysis of vendor claims, the article investigates how practitioners and experts engage with these tools and how their responses reveal tensions around interpretive agency, ethical design practice and inclusivity.

In doing so, the article contributes to emerging debates in branding, media theory and AI ethics. It argues for a reflexive and value-conscious approach to AI deployment; one that foregrounds transparency, cultural sensitivity and human-centred design. The following sections trace how algorithmic authority is constructed, negotiated and resisted across professional contexts and what this reveals about the future of branding in a hyperdigital world.

1.1. AI-supported Visual Evaluation in Hyperdigital Branding

AI has become a central actor in the design, delivery and evaluation of brand communications. Its integration into advertising and branding practices reflects a broader shift toward automation, predictive modelling and data-driven optimisation (Deloitte, 2025; McKinsey & Company, 2023; Steimer, 2019). AI-supported tools now routinely simulate consumer behaviour, generate performance metrics and offer design recommendations based on algorithmic interpretations of visual stimuli (e.g., Exner *et al.*, 2025; Harkness *et al.*, 2023; Hue & Hung, 2025). These systems are marketed as efficient, scalable and empirically grounded; qualities that appeal to practitioners operating under time and budget constraints (McKinsey & Company, 2023; Samson & Pothong, 2025). In particular, predictive attention modelling tools have gained traction for their ability to produce heatmaps, gaze plots and area-of-interest rankings without the need for human participants (Olsen, 2025). Their outputs are presented as proxies for audience perception, enabling creatives to make rapid decisions about layout, hierarchy and focal emphasis.

Yet, the epistemological foundations of these tools remain under-theorised. While they claim to replicate human attention with high fidelity – often citing benchmark comparisons with eye-tracking studies (e.g., 3M VAS, 2025; EyeQuant, 2025; Feng-GUI, 2025) – the basis for such claims is rarely transparent. The algorithms that underpin these systems are proprietary and the datasets used for training are typically undisclosed, limiting opportunities for external scrutiny and risking flattening per-

ceptual diversity and overlooking the interpretive strategies audiences use to make sense of brand content.

These limitations are particularly consequential in relation to equity, diversity and inclusion (EDI). Bias in AI systems more broadly has been well-documented, especially where training datasets underrepresent marginalised groups (Adib-Moghaddam, 2023a; Heigl, 2025; Manyika *et al.*, 2019; UNESCO IRCAI, 2024). When cultural nuance and lived experience are excluded from algorithmic modelling, the outputs may misrepresent or ignore key audience segments (Adib-Moghaddam, 2023b). While concerns about AI bias have been explored in various domains, they remain under-researched in branding and advertising communications, where recent bibliometric analyses reveal that marketing-specific investigations are sparse, fragmented and rarely address how algorithmic outputs shape visual and strategic brand decisions (Bacalhau *et al.*, 2025; Hue & Hung, 2025).

The gap in research has tangible implications for branding practice. In branding contexts, algorithmic bias can result in design recommendations that reinforce dominant norms, marginalise alternative ways of seeing and distort strategic messaging. These risks are compounded by the commercial imperatives of AI vendors, who may restrict the scope of training data to reduce development costs and maximise return on investment (Crawford, 2022; Samson & Pothong, 2025). This economic logic can prioritise efficiency over representational accuracy, further entrenching normative assumptions in brand evaluation.

By contrast, eye-tracking has long served as the methodological benchmark for studying visual attention in advertising and design (Boltz & Trommsdorff, 2022). It offers fine-grained data on gaze behaviour, such as fixation duration, saccadic movement and pupil dilation, allowing researchers to trace how viewers interact with visual content in real time. Crucially, eye-tracking supports the integration of qualitative methods, including think-aloud protocols and retrospective interviews, which link gaze data to subjective experience (Andry, 2019). This methodological depth enables researchers to understand perception as a socially and culturally situated process, rather than a purely visual or technical one. In branding, this means that design effectiveness cannot be reduced to visual prominence alone; it must also support message clarity, emotional resonance and audience relevance.

Despite its strengths, eye-tracking remains underutilised in commercial practice. High equipment costs, technical complexity and the need for trained personnel have limited its accessibility, particularly for smaller agencies and brands (Funke *et al.*, 2016; Nordfält & Ahlbom, 2024). This has created a demand for automated alternatives that promise similar insights with lower resource investment. Yet, the substitution of human-centred methods with machine-simulated outputs introduces significant risks. AI tools may replicate the surface features of eye-tracking data, such as heatmaps and gaze paths, but they can lack the interpretive nuance required to distinguish between superficial attention and meaningful engagement (Olsen, 2025). Such misattributions reveal the limitations of machine vision in accounting for lived experience. When outputs are taken at face value, they risk distorting brand messaging – by elevating visually prominent but contextually irrelevant elements – and undermining strategic intent, particularly in campaigns aimed at culturally diverse or socially sensitive audiences.

These developments raise important normative and ethical questions about authorship, agency and representation in branding practice. As AI tools become more embedded in design workflows, concerns emerge about how their outputs are interpreted and acted upon. This paper explores these dynamics through empirical engagement with practitioners and experts, examining how algorithmic authority is constructed, negotiated and resisted in professional contexts.

1.2. The Algorithmic Gaze: A Conceptual Framework

This article is grounded in the proposition that artificial intelligence does not merely assist in the evaluation of brand communications; it actively shapes the conditions under which brand perception is constructed, interpreted and acted upon. Central to this proposition is the concept of the *algorithmic gaze*, which serves as the primary analytical lens for interrogating how AI-supported visual evaluation tools simulate, encode and normalise particular ways of seeing.

The algorithmic gaze refers to the computational reproduction of human attention through machine vision systems. It draws on critical media theory and algorithm studies to conceptualise AI not as a neutral observer but as a cultural and technical actor that mediates perception (Fisher & Mehozay, 2019). Unlike the human gaze, which is shaped by embodied experience, cultural context and interpretive agency, the algorithmic gaze is governed by training data, modelling assumptions and optimisation protocols. It operationalises attention through quantifiable proxies – contrast, salience, centrality – while abstracting away from the social and cognitive dimensions of viewing. In doing so, it constructs a normative visual logic that privileges certain design elements and marginalises others, often in ways that reflect dominant cultural and commercial priorities (see Exner *et al.*, 2025).¹

Building on this definition, the algorithmic gaze can be understood not only as a mode of perception but as a mechanism of *visual governance*. Rather than passively reflecting audience behaviour, AI-supported evaluation tools intervene in the design process by shaping what is considered perceptually salient or strategically valuable. Their outputs – often presented as visual metrics – can influence design decisions and aesthetic priorities. When treated as authoritative, such metrics may reconfigure creative agency, shifting decision-making from interpretive judgement to algorithmic prescription (see Cekova *et al.*, 2025).

This also raises questions of *representation and exclusion*. Because AI systems are trained on datasets that reflect particular cultural and behavioural norms, their outputs may reproduce dominant visual logics while overlooking alternative perceptual frameworks. This has implications for how audience diversity is recognised – or marginalised – within branding environments. The framework invites scrutiny of how machine vision encodes assumptions about what is worth seeing and whose ways of seeing are legitimised (see Foka *et al.*, 2025; Numan Said *et al.*, 2025).

To critically engage with these dynamics, this article adopts a *reflexive and interpretive orientation*. AI-generated metrics are treated not as objective truths but as discursive artefacts; outputs that must be read, contextualised and interrogated. The conceptual framework foregrounds practitioner agency and supports a more situated, culturally attuned engagement with algorithmic tools. It aligns with broader calls in digital media scholarship for transparency, accountability and human-centred design

in the development and deployment of AI systems (e.g., Auernhammer, 2020; Heigl, 2025; Zhao & Xu, 2025).

The algorithmic gaze, in this sense, foregrounds the visual and perceptual dimensions of algorithmic mediation, particularly in branding environments where creative decisions are increasingly shaped by machine-simulated attention metrics. In doing so, it intersects with but remains distinct from adjacent frameworks such as ‘algorithmic governmentality’ (Rouvroy & Berns, 2013), ‘surveillance capitalism’ (Zuboff, 2019) and critical interrogations of Big Data’s epistemological claims (boyd & Crawford, 2012). These approaches have critically examined how algorithmic systems govern behaviour, produce knowledge and commodify attention. However, they tend to focus on behavioural modulation, economic extraction or epistemic automation (that is, the automation of knowledge production and decision-making). The algorithmic gaze, by contrast, centres on the governance of visual culture and interpretive agency. It extends Striphas’s (2015) notion of ‘algorithmic culture’ by showing how aesthetic judgement is increasingly automated, privatised and rendered opaque. In doing so, it critiques the displacement of interpretive agency and the normative assumptions embedded in computational seeing. It offers a framework for understanding how AI-supported tools mediate visual perception in branding, not only through technical processes but through cultural and ethical logics. This conceptual lens opens space for alternative approaches that prioritise interpretive depth, contextual knowledge and inclusive design.

While this framework is conceptually grounded, its contours are refined through empirical engagement with practitioners and experts in subsequent sections.

2. Methodological Approach and Data Sources

This study adopts a qualitative, interpretive methodology to examine the ethical implications of AI-supported visual evaluation tools in branding. Analysis across all datasets was guided by a reflexive orientation, foregrounding interpretive sensitivity and the situated nature of practitioner and expert responses. Three interrelated data sources inform the study: (1) semi-structured interviews with advertising professionals, (2) expert reflections from scholars and consultants in AI ethics and inclusive branding, and (3) promotional materials from AI tool vendors.

2.1 Practitioner Interviews: Exploring Interpretive Friction

The first dataset comprises 40 semi-structured interviews originally conducted for a comparative study exploring how AI-generated visual metrics and research-grade eye-tracking outputs inform design judgements.ⁱⁱ Participants, based in Germany, had between three and eleven years of experience in campaign ideation, planning and execution, but no prior exposure to eye-tracking technologies.

Interviews followed a two-part structure: first, participants reviewed system-generated metrics and proposed design optimisations; second, they discussed and reflected on their engagement with the metrics. A think-aloud protocol combining concurrent and retrospective elements was used to capture interpretive reasoning. Interviews were conducted via Zoom, lasted between 45 and 70 minutes and

were transcribed for analysis. All participants provided informed consent, and the study adhered to GDPR-compliant data handling and institutional ethics protocols.

These practitioner interviews were re-analysed through a reflexive lens, focusing on moments of interpretive friction; where participants questioned, resisted or contextualised algorithmic outputs. Analysis was guided by reflexive thematic analysis (Braun & Clarke, 2021), involving iterative reading, open coding, memoing and clustering codes into themes that captured ethical and epistemic tensions. Rather than applying a rigid coding frame, the process emphasised interpretive flexibility, treating responses as situated judgements about algorithmic authority in branding practice. This approach aligns with qualitative traditions that foreground reflexivity and methodological adaptability in contexts involving contested knowledge and sociotechnical complexity (*i.a.*, Trent & Cho, 2014).

2.2 Expert Interviews: Situating Practitioner Insights

The second dataset includes seven expert interviews with senior academics, analysts and consultants based in Germany and the UK. Conducted via Zoom and lasting approximately 90 minutes each, these interviews were transcribed and analysed thematically following Braun and Clarke's (2021) reflexive approach to thematic analysis. They extend the practitioner analysis by situating it within broader discourses of epistemic authority, cultural representation and ethical design. All participants provided informed consent in accordance with institutional ethics protocols. The inclusion of German experts reflects the geographic alignment with the practitioner sample, though the themes explored, such as interpretive agency and algorithmic authority, are not culturally bounded and pertain to broader dynamics of human-machine interaction.

2.3 Vendor Discourse Analysis: Constructing Algorithmic Authority

The third dataset comprises a reflexive discourse analysis of publicly available materials from 3M VAS and Feng-GUI, selected to ensure consistency with the original comparative study and alignment between practitioner experience and vendor claims. These texts were examined for rhetorical strategies and normative assumptions about perception, attention and design. Particular attention was given to how vendor discourse frames algorithmic outputs as objective and prescriptive, shaping practitioner expectations and influencing design decision-making.

2.4 Ethical Frameworks and Interpretive Orientation

The interpretive approach is informed by three complementary ethical frameworks: Floridi's concept of the infosphere (2014), feminist data ethics (D'Ignazio & Klein, 2020) and EDI-informed critiques of algorithmic bias (Benjamin, 2019; Noble, 2018). Together, these frameworks support a multi-dimensional reading of the data, sensitising the analysis to potential risks such as epistemic injustice, representational exclusion and the displacement of human interpretive agency.

Rather than evaluating technical performance, the methodology interrogates the cultural and ethical implications of AI in shaping brand perception and decision-making in hyperdigital environments. Insights from the three data sources inform the development of the algorithmic gaze framework and the Interpretive Stewardship heuristic – a model for ethical engagement with AI in branding.

3. Practitioner Perspectives: Authority and Ethical Friction

The re-analysis of practitioner interviews reveals a layered and often conflicted relationship between advertising professionals and AI-supported visual evaluation tools. While these systems are marketed as efficient proxies for human attention, their outputs are frequently met with scepticism, interpretive resistance and ethical concern. The following sections are organised around practitioner-identified themes that emerged through reflexive re-analysis. While these themes resonate with the conceptual dimensions of the algorithmic gaze, they were not used as a coding framework and reflect situated professional concerns.

3.1. Navigating Algorithmic Design Authority

A recurring theme is the initial plausibility of AI-generated metrics, particularly heatmaps and gaze plots, which often resemble those produced by research-grade eye trackers. Practitioners described these outputs as *“visually convincing”* (P21), with one noting, *“everything that is important gets a decent amount of attention”* (P22). However, this confidence often eroded under closer scrutiny. AI systems frequently flagged elements such as disclaimer text, typically ignored in real-world viewing, as high-attention zones. This led to widespread scepticism: *“This must be an error”* (P31/P32/P36); *“That’s simply not how people would view this advert in normal life”* (P25).

These responses reveal a deeper concern: the disconnect between algorithmic outputs and culturally conditioned viewing practices. The algorithmic gaze, as conceptualised in this article, operates through computational proxies – contrast, salience, centrality – that abstract away from the lived experience of perception. In privileging technically prominent elements, it risks misrepresenting what is socially meaningful. The repeated misattribution of attention to disclaimers or decorative motifs exemplifies this logic. While such elements may be visually salient, they are often culturally conditioned to be ignored. The algorithm’s failure to recognise this distinction exposes its blind spot: it sees what is technically visible, not what matters in context.

This dynamic reflects a broader concern about delegated decision-making in hyperdigital branding environments. As AI tools become more embedded in creative workflows, there is a risk that practitioners will increasingly rely on algorithmic outputs to guide creative choices. This reliance may be driven by time pressures, managerial expectations or the perceived objectivity of data. As one participant noted, *“There’s pressure to act on the data, even when it doesn’t make sense. If the dashboard says the logo isn’t working, someone will want it changed, regardless of context”* (P33). Another added, *“I’ve seen teams defer to the tech-facilitated output just to avoid debate. It becomes a shortcut for decision-making, but not always a good one. Particular younger colleagues who don’t want to go heads-on with others within the team”* (P39).

Despite the interpretive limitations and contextual misalignments, practitioners often found AI-generated outputs compelling, largely because of their visual clarity and ease of use. Practitioners consistently praised AI-supported tools for their accessible visualisations, describing them as

“straightforward” (P27), *“easy to use”* (P40) and *“self-explanatory”* (P25). In contrast, research-grade eye-tracking metrics were often described as *“confusing”* (P3), *“overwhelming”* (P19) and *“intimidating”* (P8), which occasionally made them more difficult to adopt, even when they offered greater methodological rigour. This usability gap reflects a deeper epistemological shift in branding environments; one that can be understood through Floridi’s (2014) concept of the infosphere. AI-supported tools do not merely assist practitioners; they actively shape the informational conditions under which decisions are made. By determining what is rendered visible, relevant or actionable, these systems reconfigure the boundaries of perception and judgement.

3.2. Cultural Blind Spots and Inclusive Design Challenges

Practitioners also expressed concern about the opacity of the systems. Many noted that they were unable to interrogate how attention was modelled or what assumptions were embedded in the metrics. As one participant put it, *“I don’t know how the software decides what’s important. There’s no way to question the logic behind the heatmap. (...) It remains a mystery”* (P37). Another added, *“You’re shown a gaze plot and expected to trust it, but there’s no transparency. It’s hard to challenge something when you don’t know how it was built”* (P23). These interpretive challenges are amplified by the rhetorical strategies used in vendor marketing materials, which shape practitioner expectations and influence how outputs are received and acted upon. One AI-supported visual evaluation tool describes its system as *“like having a visual spellcheck”*, claiming *“92% accuracy”* and encouraging users to create *“designs with impact”* (3M VAS, 2025). Similar rhetorical strategies are found with other AI tools, with another asserting that its reports *“closely resemble a 5 seconds eye tracking session of 40 people”* and that its algorithms are *“MIT approved”* (Feng-GUI, 2025). These claims position the tools as authoritative, predictive and scientifically validated, despite the lack of high-quality, peer-reviewed evidence. The language used – *“the science behind”*, *“data-driven decisions”*, *“AI-powered neuromarketing”*, *“precision”*, *“confidence”* – frames the outputs as definitive rather than interpretive, reinforcing the prescriptive nature of the algorithmic gaze.

This prescriptive framing not only shapes practitioner behaviour but also aligns with broader critiques of datafication, where complex human perception is reduced to quantifiable proxies. These concerns about interpretive authority and practitioner vulnerability are further illuminated by feminist data ethics, which offer a critical lens on the structural limitations of AI systems, particularly their lack of transparency and cultural sensitivity. Building on earlier critiques of epistemic asymmetry in data-driven decision-making – where access to or control over knowledge is unevenly distributed – feminist perspectives emphasise the importance of situated knowledge, contextual sensitivity and the recognition of diverse ways of seeing. The reduction of perceptual experience to visual metrics such as salience scores and fixation paths strips away the nuance of audience engagement and narrows the scope of design evaluation. As one practitioner noted, *“It feels like the software is telling me what people should look at, not what they actually do. There’s a kind of design bias built into these metrics”* (P40).

The implications for equity, diversity and inclusion are especially noteworthy. If AI-supported tools are trained on datasets that reflect dominant cultural norms, they may reproduce exclusionary design

practices and overlook the needs of marginalised audiences. As one practitioner observed, *“The system doesn’t seem to recognise that different audiences respond differently. It treats all viewers as the same, which is a problem for inclusive design”* (P32). Another added, *“We work with multicultural campaigns and these datasets [from the AI tools] seem often to miss culturally specific cues. It’s not just inaccurate – it’s exclusionary”* (P34). Without cultural calibration, AI metrics risk misrepresenting audience behaviour and reinforcing visual hierarchies that privilege certain groups over others.

3.3. Interpretive Agency and Ethical Engagement

Several practitioners were acutely aware of these limitations. Their responses demonstrate a form of interpretive agency that resists the authority of the algorithm, particularly among senior professionals. These individuals often drew on their experience to critically assess the outputs, contextualise anomalies and override recommendations that conflicted with strategic intent. One senior practitioner remarked: *“The metrics flagged the product name as the main attention hotspot, but that’s not how people read this kind of layout. It’s central, yes, but it’s not emotionally or narratively significant. I’d rather trust what I know about how people scan for meaning than what this heatmap says”* (P35). Another senior strategist noted: *“The [AI] system said the call-to-action was underperforming, but it’s placed exactly where it needs to be for this demographic. I think it doesn’t understand the context here”* (P24). These responses exemplify a mode of engagement that treats AI metrics not as definitive answers but as interpretive prompts, subject to scrutiny, contextualisation and, where necessary, rejection.

By contrast, junior practitioners exhibited greater uncertainty. When confronted with conflicting or ambiguous metrics, they were more likely to defer to the data or reject it wholesale, lacking the confidence to navigate its interpretive complexity. One junior advertiser explained: *“I wasn’t sure what to make of the gaze plot. It looked convincing, but I didn’t know if it was right. I just went with it because it seemed safer than arguing against it when I have no other data to say differently”* (P27). This quote reveals a layered ethical tension. The practitioner recognises the persuasive power of the visualisation but lacks the epistemic tools to challenge it. The decision to ‘go with it’ reflects not just a lack of confidence, but a structural absence of alternative evidence or interpretive support.

Finally, the analysis underscores the importance of reflexive practice. Practitioners who engaged critically with the metrics – questioning anomalies, contextualising outputs and integrating their own experience – were better able to navigate the limitations of the tools. Their responses exemplify a mode of engagement that resists the passive consumption of data and affirms the value of human judgement. This stance aligns with the ethical imperative of maintaining interpretive agency in the face of algorithmic authority.

4. Expert Perspectives: Strategic, Ethical and Conceptual Insights

While practitioner interviews revealed interpretive friction and epistemic vulnerability in everyday use, expert reflections extend this critique by interrogating the strategic, ethical and conceptual foundations of AI-supported evaluation tools.

4.1. Strategic and Normative Dimensions of the Algorithmic Gaze

The concept of the algorithmic gaze, introduced in this article resonated strongly with experts. One contributor, a senior academic specialising in AI and advanced computing, described it as *“capturing something crucial: these systems don’t just observe – they encode assumptions about what matters visually, based on training data and model architecture.”* This comment underscores that predictive attention systems are not passive tools but active agents in shaping visual culture. Another senior academic specialising in gender and technology, connected the concept to feminist critiques of vision and representation. *“It’s not just about what is seen, but about who gets to define what is worth seeing”,* they noted. *“These systems risk reinforcing dominant visual logics that marginalise alternative ways of looking.”* This framing situates the algorithmic gaze within a lineage of critical theory that interrogates the politics of perception, suggesting that AI tools may reproduce the same exclusions and hierarchies long critiqued in media and design studies.

A third contributor, working in human-computer interaction, described the concept as *“a useful provocation. It helps us ask not just whether the metrics are accurate, but what kind of visual culture they are producing.”* This shift, from performance to cultural production, marks a significant departure from conventional evaluations of AI tools, which tend to focus on technical fidelity rather than normative influence.

The normative and political implications of this shift were explored through the lens of Floridi’s infosphere (2014). One expert argued that *“these tools don’t just analyse visual content – they shape the conditions under which design decisions are made. They influence what is considered legible, effective and valuable, often without users being aware of the assumptions behind the metrics.”* This observation aligns with the article’s earlier critique of vendor discourse, where AI-supported software solutions are described as *“developed by expert neuro- and data scientists”* (3M VAS, 2025) and claim *“MIT-approved”* accuracy (Feng-GUI, 2025). Such language constructs a sense of scientific legitimacy that obscures the normative nature of the outputs.

An digital ethicist affiliated with a multinational think tank expanded on this point, stating: *“The problem is not just technical opacity – it’s epistemic authority. These systems present themselves as objective, but they are deeply normative. They tell creatives what to prioritise, often based on unexamined assumptions about attention, cognition and desirability.”* This critique underscores the need to interrogate not only the outputs but the rhetorical strategies through which AI tools gain legitimacy. The expert added that *“these tools are shaping the information environment of advertising and branding – what counts as evidence, what gets seen and what gets ignored.”*

Transparency emerged as a central concern across all expert responses. One contributor noted, *“You’re shown an end product; that’s all. The opacity of these systems undermines accountability. It’s a black box and that’s a problem for both researchers and practitioners.”* Another expert argued that *“without clear documentation of model architecture, training data and validation methods, it’s impossible to assess the reliability or fairness of the outputs.”* A third framed the issue in terms of feminist data ethics: *“Opacity is not just a technical flaw – it’s a political one. It prevents users from understanding how their decisions are being shaped and it reinforces asymmetries of power between tool developers and tool users.”*

One digital ethicist proposed that vendors should be required to publish model documentation and bias audits, similar to emerging standards in algorithmic accountability. *“To guide strategic decisions effectively, these tools must meet a higher standard of transparency”*, they argued. This recommendation aligns with broader calls in AI ethics for explainability, auditability and participatory governance (e.g., Auernhammer, 2020; Heigl, 2025; Zhao & Xu, 2025). It also resonates with recent developments in UK advertising and policy, including the Advertising Standards Authority’s 2024-2028 strategy and *Future-proof Advertising: How will AI change advertising and regulation?* (Advertising Standards Authority, 2023, 2025), which emphasises the need for more effective, responsive regulation regarding AI in branding practice. The Advertising Standards Authority’s proactive monitoring of AI activity within the industry, alongside the UK Government’s strategies and white papers (Brione & Gajjar, 2024; DSIT, 2025), signals a growing institutional commitment to responsible AI governance in commercial communications.

4.2. Cultural Representation, Bias and Inclusive Design

The implications for equity, diversity and inclusion were explored in depth by three senior brand consultants specialising in technology use and inclusive communications. One consultant observed, *“These systems often assume a universal viewer, but in reality, perception is shaped by culture, context and lived experience. If the model doesn’t account for that, it risks excluding entire audience segments.”* Another reflected on the broader strategic risks posed by algorithmic evaluation, noting that *“AI tools often flatten cultural nuance into generic metrics. This can lead to misalignment between brand messaging and audience values, especially in campaigns designed for diverse or marginalised communities.”* The third consultant echoed this concern and further emphasised the need for cultural calibration: *“For use in inclusive branding, these tools must be trained on diverse datasets and tested across demographic groups. Otherwise, they risk reinforcing dominant norms and undermining the goals of inclusive design.”* All three consultants also raised the importance of transparency in terms of the trained data, which would empower users to critically engage with the findings better.

These reflections extend the critique of the algorithmic gaze by showing how it can conflict with EDI principles. They also suggest concrete directions for tool improvement, including dataset diversification, participatory testing and the integration of cultural metadata into attention models. The experts were clear that inclusive design cannot be automated without context: *“You can’t outsource cultural sensitivity to an algorithm”*, one stated. *“It has to be built into the process, not bolted on afterwards.”*

4.3. Interpretive Agency and Ethical Accountability

The question of human judgement and interpretive agency was another recurring theme. One expert warned that *“when metrics are treated as definitive, they displace professional judgement. That’s dangerous; not just ethically, but strategically.”* Another added, *“Design is not just a technical exercise, it’s a cultural and communicative one. Reducing it to salience scores ignores its complexity and undermines its social function.”*

The brand experts proposed that AI tools should be designed to support interpretive dialogue rather than prescriptive output. *“Instead of telling users what to do, the system could present multiple interpretations, highlight uncertainty and invite reflection. That would be a more ethical and empowering approach.”* Others added: *“We must use the metrics as a starting point, not a conclusion. They’re useful, but only when combined with experience, intuition and audience insight.”*; *“There’s a risk that junior teams will defer to the data too quickly. We’ve had to train people to treat the outputs critically, not as gospel.”*

Across these reflections, a shared set of strategic priorities emerged: the need for transparent systems, inclusive design processes and critical literacy around ethics and power. One contributor summarised the strategic imperative: *“If these tools are going to be part of the design process, they need to be accountable, interpretable and inclusive. Otherwise, they risk doing more harm than good.”*

5. Synthesis of Practitioner and Expert Insights

Practitioners and experts alike recognise that AI-supported visual evaluation tools do more than measure attention; they actively shape how visual meaning is constructed and legitimised. While practitioners engage with these systems in the immediacy of creative workflows, experts reflect on their broader epistemic and strategic consequences. Together, these insights shine a light on how the algorithmic gaze operates not only as a technical function but as a mode of visual governance – prescribing design priorities, configuring perception and redefining the conditions of creative agency.

Practitioner responses reveal a persistent tension between operational utility and interpretive trust. Senior professionals often resisted algorithmic outputs when they conflicted with contextual knowledge or strategic goals, demonstrating reflexive judgement. Junior practitioners, by contrast, tended to defer to the metrics, particularly in the absence of alternative evidence or institutional support. This asymmetry underscores a deeper issue: the uneven distribution of interpretive agency in hyperdigital branding environments. Where algorithmic outputs are treated as authoritative, those with less experience or confidence may feel compelled to follow them, even when they contradict lived experience or communicative intent. Over time, this dynamic risks institutionalising algorithmic authority; not only in creative decisions but in the culture of branding itself, where data-driven outputs are privileged over human insight and critical reflection.

Expert reflections extend these practitioner concerns, situating them within broader discourses of epistemic authority, normative design logic and ethical accountability. They affirm the algorithmic gaze as a critical tool for understanding how AI systems encode assumptions about attention, value and

visibility. Importantly, they highlight that the risks posed by these systems – exclusion, misrepresentation, homogenisation – are not incidental but structurally embedded in how algorithmic outputs are operationalised and legitimised.

To move beyond critique, the synthesis draws on Floridi's (2014) concept of the infosphere to re-frame branding environments as informational ecosystems shaped by algorithmic systems and vendor discourse. These systems act as informational agents, determining what counts as evidence and how creative agency is distributed. This raises ethical questions about the design of the infosphere itself: who controls its architecture, whose values are encoded and how transparency and accountability are maintained. Reframing branding practice through Floridi's lens invites a shift from tool adoption to infosphere stewardship; a responsibility to cultivate environments that support ethical reflection, pluralistic knowledge and inclusive design.

Building on this reframing, feminist data ethics extend the synthesis by shifting the focus from critique to co-liberation (D'Ignazio & Klein, 2020). Co-liberation reframes ethical design not as harm mitigation alone, but as a collective project of empowerment; enabling diverse practitioners to challenge normative systems and reshape design culture. Inclusive design literacy becomes not just a technical skill but a political strategy, equipping practitioners to resist epistemic asymmetry and assert interpretive authority. Embedding co-liberation into branding workflows means designing tools that support dialogue, uncertainty and multiplicity, not just optimisation. It also means recognising that ethical design is inseparable from structural change: cultivating cultures of critique, redistributing interpretive power and ensuring that AI systems serve the goals of equity, not just efficiency.

What emerges from the synthesis is a sharpened conceptual focus on the algorithmic gaze as a mechanism through which normative assumptions about perception, value and design are encoded and enacted. AI-supported evaluation tools must be understood not as neutral instruments of optimisation, but as cultural artefacts that require critical negotiation. Both practitioners and experts advocate a shift from passive adoption to active interrogation, foregrounding human judgement, contextual nuance and ethical responsibility. This reorientation of branding practice includes cultivating design literacy, embedding reflexive protocols into creative workflows and ensuring that algorithmic systems support rather than supplant creative and strategic agency. It also demands transparency, inclusivity and interpretive depth; anchored in ethical stewardship of the infosphere and guided by co-liberative design principles.

To consolidate these insights into a usable framework, this article introduces a heuristic for evaluating AI-supported visual tools in branding environments: the *Interpretive Stewardship* (see Figure 1). Developed in response to the normative and epistemic risks surfaced through the lens of the algorithmic gaze, this model offers a scaffold for resisting these tendencies and cultivating inclusive, reflexive creative practice. Interpretive Stewardship comprises three interlinked dimensions, each addressing a distinct facet of ethical branding in hyperdigital environments:

1. *Epistemic Transparency*: Drawing from Floridi's (2014) infosphere, this dimension emphasises the need for visibility into model architecture, training data and decision logic. It reframes AI tools as informational agents whose influence must be made legible and accountable.
2. *Situated Interpretation*: Informed by feminist data ethics and EDI-informed critique, this dimension foregrounds the importance of contextual knowledge, lived experience and cultural nuance in interpreting algorithmic outputs. It resists the flattening of perception into abstract metrics and supports pluralistic ways of seeing.
3. *Co-Liberative Practice*: Extending D'Ignazio and Klein's (2020) notion of co-liberation, this dimension calls for participatory design, inclusive testing and critical literacy. It positions ethical engagement not as a reactive safeguard but as a proactive strategy for redistributing interpretive power.

Together, these dimensions offer a scaffold for reflexive engagement with AI tools – one that supports human judgement, resists normative closure and cultivates ethical stewardship of the branding infosphere.

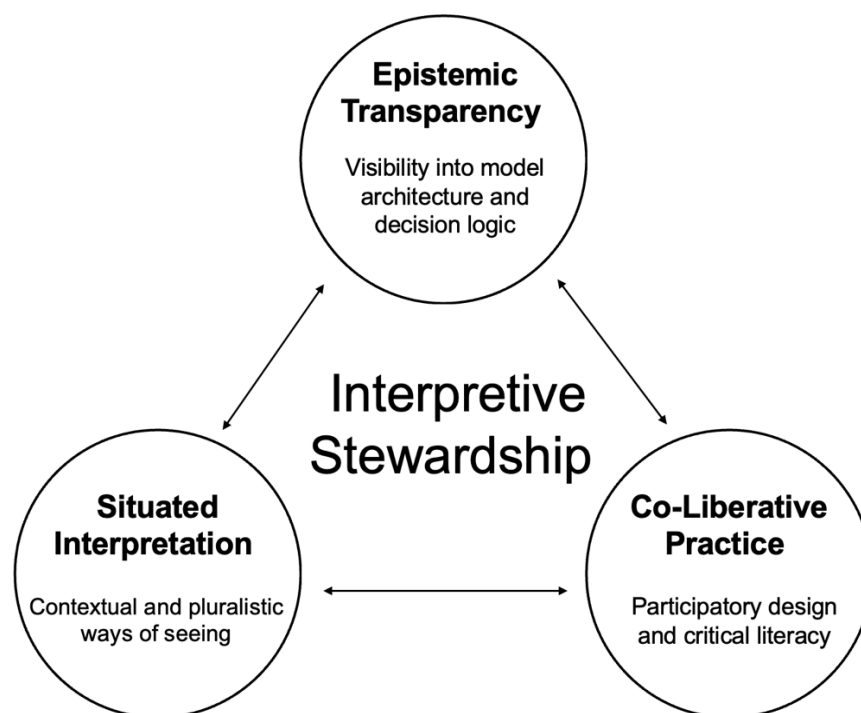


Figure 1. Interpretive Stewardship: A Heuristic for Ethical Engagement with AI in Branding

To operationalise the Interpretive Stewardship heuristic, branding teams must adopt three inter-linked practices. First, epistemic transparency should be a baseline requirement: only AI tools with clear, auditable documentation of model architecture, training data and decision logic should be considered for implementation. Second, situated interpretation must be actively supported. Practitioners should be encouraged, and institutionally enabled, to contextualise algorithmic outputs using audience insight, cultural nuance and strategic intent. This includes embedding interpretive protocols into design workflows and resisting reductive metrics. Third, co-liberative practice requires targeted training

in inclusive design literacy, enabling teams to critically engage with AI outputs and challenge normative assumptions. Guidelines should explicitly promote this reflexive engagement, ensuring that automation augments rather than overrides human judgement.

6. Conclusion

AI-supported visual evaluation tools are reshaping the epistemic and ethical foundations of branding. This article has shown that these systems do not simply simulate human attention; they prescribe visual hierarchies that reflect modelling assumptions, cultural biases and commercial imperatives. In doing so, they reconfigure how brand meaning is constructed, interpreted and legitimised in hyperdigital environments.

The concept of the *algorithmic gaze* developed here offers a critical lens for understanding this transformation. It bridges branding, media theory and AI ethics, enabling a shift from evaluating AI tools based on performance metrics to interrogating their cultural and normative influence. By applying ethical reflection frameworks – Floridi's infosphere, feminist data ethics and EDI-informed critiques – the article reframes AI metrics as discursive artefacts that risk narrowing interpretive possibilities and embedding design biases that exclude culturally specific ways of seeing. Methodologically, the combination of practitioner and expert insights and vendor discourse critique exposes how algorithmic authority is constructed and sustained through rhetorical and technical opacity.

These contributions are significant for both scholars and practitioners. For researchers, the article offers a vocabulary and analytical toolkit for studying AI-human interaction in branding. For practitioners, it highlights the need to maintain interpretive agency when working with algorithmic outputs, ensuring that data-informed decisions remain aligned with strategic goals, audience insight and creative intent.

In response to this article's research question, the findings demonstrate that AI-generated visual metrics influence brand communications not only by guiding creative decisions, but by shaping the conditions under which those decisions are made. The risks of epistemic distortion, representational exclusion and the erosion interpretative agency are not peripheral but central to how these tools operate.

To address these risks, the article proposes the *Interpretive Stewardship* heuristic – a three-dimensional framework comprising epistemic transparency, situated interpretation and co-liberative practice. This model offers a scaffold for reflexive engagement with AI tools in branding environments, supporting human judgement, resisting normative closure and cultivating ethical stewardship of the infosphere. Future research could operationalise this heuristic through participatory audits, cross-cultural testing protocols and longitudinal studies of practitioner engagement. Such approaches would help evaluate how interpretive agency is distributed, how ethical awareness evolves and how branding teams negotiate the balance between automation, creativity and inclusivity.

Institutionally, these findings align with emerging regulatory priorities, such as the UK Advertising Standards Authority's 2024-2028 strategy and the recent calls by the UK Government for rapid adoption of AI across sectors. Embedding human-centred values into AI systems, through transparency in modelling assumptions, cultural calibration of training data and participatory development processes, is not only an ethical imperative but increasingly a regulatory and socio-political expectation.

As the algorithmic gaze continues to shape the visual culture of branding, the challenge is not only to refine the tools, but to interrogate and reconfigure the values – about perception, authority and design – that these systems encode. Ensuring that branding remains a site of cultural negotiation rather than computational prescription will require a sustained commitment to co-liberation, ethical reflexivity and inclusive design literacy.

References

- 3M VAS. (2025). *Fast data while you design*. <https://vas.3m.com/>
- Adib-Moghaddam, A. (2023a). *Is Artificial Intelligence Racist? The Ethics of AI and the Future of Humanity*. Bloomsbury Academic.
- Adib-Moghaddam, A. (2023b). For minorities, biased AI algorithms can damage almost every part of life. *The Conversation*. <https://doi.org/10.64628/AB.4KEN7X5HS>
- Advertising Standards Authority. (2023, November 28). *AI-assisted collective ad regulation. The ASA's 2024-2028 strategy*. <https://www.asa.org.uk/news/ai-assisted-collective-ad-regulation-our-new-strategy.html>
- Advertising Standards Authority. (2025, February 7). *AI, advertising, and the policy landscape – CAP proactive monitoring*. <https://www.asa.org.uk/news/ai-advertising-and-the-policy-landscape-cap-proactive-monitoring.html>
- Andry, T. (2019). The Qualitative Analysis in Eye Tracking Studies: Including Subjective Data Collection in an Experimental Protocol. *Human Interface and the Management of Information. Information in Intelligent Systems. HCII 2019. Lecture Notes in Computer Science, 11570 LNCS*, 335–346. https://doi.org/10.1007/978-3-030-22649-7_27
- Auernhammer, J. (2020). Human-centered AI: The role of Human-centered Design Research in the development of AI. *DRS Biennial Conference Series. Synergy, 11-14 August*, 1315–1333. <https://doi.org/10.21606/DRS.2020.282>
- Bacalhau, L. M., Pereira, M. C., & Neves, J. (2025). A bibliometric analysis of AI bias in marketing: field evolution and future research agenda. *Journal of Marketing Analytics*, 13(2), 308–327. <https://doi.org/10.1057/S41270-025-00379-6>
- Benjamin, R. (2019). *Race After Technology: Abolitionist Tools for the New Jim Code*. Polity Press.
- Boltz, D.-M., & Trommsdorff, V. (2022). *Konsumentenverhalten*. Kohlhammer. <https://doi.org/https://doi.org/10.17433/978-3-17-037789-9>
- boyd, d., & Crawford, K. (2012). Critical Questions for Big Data: Provocations for a Cultural, Technological, and Scholarly Phenomenon. *Information, Communication & Society*, 15(5), 662–679. <https://doi.org/10.1080/1369118X.2012.678878>
- Braun, V., & Clarke, V. (2021). *Thematic Analysis: A Practical Guide*. SAGE.
- Brione, P., & Gajjar, D. (2024). *Artificial intelligence: ethics, governance and regulation*. <https://doi.org/10.58248/HS51>
- Cekova, D., Corsetti, L., Ferretti, S., & Vaca, S. (2025). *Considerations and Practical Applications for Using Artificial Intelligence (AI) in Evaluations*.

<https://cgspace.cgiar.org/server/api/core/bitstreams/ea38811e-c2b6-4f31-907b-aa546c26ff6b/content>

- Crawford, K. (2022). *Atlas of AI. Power, Politics, and the Planetary Costs of Artificial Intelligence*. Yale University Press.
- Deloitte. (2025). *Now decides next: Generating a new future. Deloitte's State of Generative AI in the Enterprise*. <https://www.deloitte.com/content/dam/assets-shared/docs/about/2025/quarter-4.pdf>
- D'Ignazio, C., & Klein, L. F. (2020). *Data Feminism*. MIT Press.
- DSIT - Department for Science, Innovation & Technology. (2025). *AI Opportunities Action Plan*. <https://www.gov.uk/government/publications/ai-opportunities-action-plan/ai-opportunities-action-plan>
- Erickson, K. (2024). AI and work in the creative industries: digital continuity or discontinuity? *Creative Industries Journal*, 1–21. <https://doi.org/10.1080/17510694.2024.2421135>
- Exner, Y., Hartmann, J., Netzer, O., & Zhang, S. (2025). *AI in Disguise – How AI-Generated Ads' Visual Cues Shape Consumer Perception and Performance* (Working Paper). <https://doi.org/10.2139/SSRN.5096969>
- EyeQuant. (2025). *Data Driven Design. The Future Is Here*. <https://www.eyequant.com/>
- Feng-GUI. (2025). *AI-Powered Visual Analytics*. <https://feng-gui.com/>
- Fisher, E., & Mehozay, Y. (2019). How algorithms see their audience: media epistemes and the changing conception of the individual. *Media, Culture & Society*, 41(8), 1176–1191. <https://doi.org/10.1177/0163443719831598>
- Floridi, L. (2014). *The Fourth Revolution. How the Infosphere Is Reshaping Human Reality*. Oxford University Press.
- Foka, A., Griffin, G., Ortiz Pablo, D., Rajkowska, P., & Badri, S. (2025). Tracing the bias loop: AI, cultural heritage and bias-mitigating in practice. *AI & SOCIETY*. <https://doi.org/10.1007/S00146-025-02349-Z>
- Funke, G., Greenlee, E., Carter, M., Dukes, A., Brown, R., & Menke, L. (2016). Which Eye Tracker Is Right for Your Research? Performance Evaluation of Several Cost Variant Eye Trackers. *Proceedings of the Human Factors and Ergonomics Society*, 60(1), 1240–1244. <https://doi.org/https://doi.org/10.1177/1541931213601289>
- Harkness, L., Robinson, K., Stein, E., & Wu, W. (2023). *How generative AI can boost consumer marketing*. <https://www.mckinsey.com/capabilities/growth-marketing-and-sales/our-insights/how-generative-ai-can-boost-consumer-marketing>
- Heigl, R. (2025). Generative artificial intelligence in creative contexts: a systematic review and future research agenda. *Management Review Quarterly*. <https://doi.org/10.1007/s11301-025-00494-9>
- Hue, T. T., & Hung, T. H. (2025). Impact of artificial intelligence on branding: a bibliometric review and future research directions. *Humanities and Social Sciences Communications*, 12, 1–11. <https://doi.org/10.1057/s41599-025-04488-6>
- Manyika, J., Silberg, J., & Presten, B. (2019). What Do We Do About the Biases in AI? *Harvard Business Review*. <https://hbr.org/2019/10/what-do-we-do-about-the-biases-in-ai>
- McKinsey & Company. (2023). *AI-powered marketing and sales reach new heights with generative AI*. <https://www.mckinsey.com/capabilities/growth-marketing-and-sales/our-insights/ai-powered-marketing-and-sales-reach-new-heights-with-generative-ai>
- Noble, S. U. (2018). *Algorithms of Oppression: How Search Engines Reinforce Racism*. NYU Press.
- Nordfält, J., & Ahlbom, C. P. (2024). Utilising eye-tracking data in retailing field research: A practical guide. *Journal of Retailing*, 100(1), 148–160. <https://doi.org/10.1016/J.JRETAI.2024.02.005>
- Numan Said, M., Zaidi, A., Usman, R., Okon, S., Medepalli, P., Zhu, K., Sharma, V., & O'Brien, S. (2025). Deconstructing Bias: A Multifaceted Framework for Diagnosing Cultural and Compositional Inequities in Text-to-Image Generative Models. *ICLR 2025 Workshop on*

Synthetic Data × Data Access Problem.

<https://doi.org/https://doi.org/10.48550/arXiv.2505.01430>

- Olsen, D. A. (2025). AI Meets Advertising: Evaluating AI-Supported Automated Eye-Tracking Solutions to Optimise Advertising Campaigns. In V. Allemann-Ravi (Ed.), *New Generation Communication: Die Kommunikation in einer veränderten Welt* (pp. 215–241). Springer. https://doi.org/10.1007/978-3-658-47591-8_8
- Rouvroy, A., & Berns, T. (2013). Algorithmic governmentality and prospects of emancipation. Disparateness as a precondition for individuation through relationships? *Réseaux*, 177(1), 163–196. <https://doi.org/10.3917/RES.177.0163>
- Samson, R., & Pothong, K. (2025). *A learning curve? A landscape review of AI and education in the UK*. (Discussion Paper). <https://www.adalovelaceinstitute.org/wp-content/uploads/2025/01/Ada-Lovelace-Institute-Nuffield-Foundation-A-learning-curve.pdf>
- Steimer, S. (2019). The CMO Survey. August 2019 CMO Survey: Hiring, AI on the Rise. In *American Marketing Association*. <https://www.ama.org/marketing-news/august-2019-cmo-survey-hiring-ai-on-the-rise/>
- Striphas, T. (2015). Algorithmic culture. *European Journal of Cultural Studies*, 18(4–5), 395–412. <https://doi.org/https://doi.org/10.1177/1367549415577392>
- Trent, A., & Cho, J. (2014). Interpretation Strategies: Appropriate Concepts. In P. Leavy (Ed.), *The Oxford Handbook of Qualitative Research* (pp. 639–657). Oxford University Press. <https://doi.org/10.1093/OXFORDHB/9780199811755.013.021>
- UNESCO IRCAI. (2024). *Challenging Systematic Prejudices. An Investigation into Gender Bias in Large Language Models*. <https://unesdoc.unesco.org/ark:/48223/pf0000388971>
- Zhao, C., & Xu, W. (2025). Human-AI Interaction Design Standards. In W. Xu (Ed.), *Handbook of Human-Centered Artificial Intelligence*. Springer. <https://arxiv.org/abs/2503.16472>
- Zuboff, S. (2019). *The Age of Surveillance Capitalism. The Fight for a Human Future at the New Frontier of Power*. Profile Books.

ⁱ Note: The conceptual framework of the algorithmic gaze draws on a synthesis of insights from recent scholarship in AI ethics, media theory and advertising research. While the cited works do not directly theorise the algorithmic gaze, their analyses inform the interpretive logic developed here.

ⁱⁱ Olsen (2025) employed a two-phase design. In the first phase, eight static advertisements, selected for varying levels of design complexity, were evaluated using two research-grade eye-tracking systems (SMI/RED and Tobii T60) and two AI-supported software tools (3M VAS and Feng-GUI). The objective was to generate comparable visual metrics (e.g., heatmaps, gaze plots, focus maps, AOI rankings) across systems to prepare stimulus material for practitioner interviews. Due to the limited transparency of the AI tools, only headline outputs were analysed, focusing on the interpretive value of the metrics rather than their statistical equivalence. In the second phase, 40 advertising professionals based in Germany participated in semi-structured interviews. Ten participants were assigned to each system, with a mix of junior and senior professionals, all of whom had experience in campaign ideation, planning and execution but no prior exposure to eye-tracking technologies. Interviews followed a think-aloud protocol and explored how practitioners interpreted the metrics and formulated design recommendations. The current article re-analyses these interviews through an ethical and reflexive lens.