

Characteristics of AI-generated Images

AUTHOR

Nicola Döring, Thuy Dung Le, Anastasiia Shevtsova

KEYWORDS

Generative artificial intelligence (genAI), AI-generated images, synthetic images, synthetic media, text-to-image (TTI, T2I) generators

ABSTRACT

News and other non-fiction media – including educational material, advertising, science communication, and social media infotainment – increasingly use informational images generated with artificial intelligence (AI) tools. AI-generated images are visual outputs produced by computational models, most commonly diffusion models or generative adversarial networks, that synthesize novel imagery from learned patterns in large training datasets instead of capturing any real-world scene or object through a camera (Sanseviero et al., 2024). Unlike traditional digital editing or illustration, AI image generators can produce photorealistic images or stylistically coherent visuals (e.g., illustrations, diagrams) from simple text prompts. While AI-generated images have the potential for meaningful and helpful visualizations, they are also criticized – for example, for their lack of accuracy and diversity. These concerns are particularly salient in informational contexts, where images are expected to represent reality and where misleading or skewed visuals can shape audience beliefs, reinforce stereotypes, and erode trust in media. Drawing on theories about AI image generation, AI content moderation, AI generation errors, and AI biases (mainly representational and presentational biases), this entry introduces a codebook with a set of four key concepts and related variables and operational definitions designed to allow coders to consistently identify potentially problematic characteristics of AI-generated informational images.

FIELD OF APPLICATION/THEORETICAL FOUNDATION

In the field of visual communication (Thelander, 2018), images generated with artificial intelligence (AI) tools have gained increasing relevance for researchers and practitioners. AI image generators are now easily accessible through public text-to-image tools (TTI or T2I generators), leading to the growing use of AI-generated images across a broad range of entertainment (e.g., AI-generated pornography: Lapointe et al., 2025) and informational contexts (e.g., AI-supported health communication: Döring & Shevtsova, 2026). This widespread availability and adoption have made the characteristics and effects of AI-generated images an important topic of communication research.

When used for informational visual communication, AI-generated images offer genuine opportunities: they may help explain complex structures or processes, attract attention, foster reflection or improve the persuasiveness of messages (e.g., Amri & Hisan, 2023; Reed et al., 2023). At the same time, there is growing concern that the broader use of AI image generators introduces new risks, including the dissemination of large amounts of low-quality AI images (so-called AI slop) that elicits fatigue and distrust in media, the deliberate generation of AI disinformation, and the unintentional creation and dissemination of problematic AI images by well-meaning communicators due to the technological limitations of AI models (e.g., so-called AI hallucinations or confabulations; Buzzaccarini et al., 2024).

Developers of AI image generators are aware of the last-mentioned risk. They take different measures to avoid problematic AI-generated outputs and, hence, aim to design responsible AI (Gunasekara et al., 2025): they curate and filter training data, train their models to produce prompt-aligned and ethics-conforming images, and implement content moderation to block outputs that could violate user policies or laws (Döring & Shevtsova, 2026). However, it remains



<https://doi.org/10.34778/927a3q69>

© 2026, the authors. This work is licensed under the “Creative Commons Attribution – NonCommercial – NoDerivatives 4.0 International” license (CC BY-NC-ND 4.0)



an empirical question how successful different AI companies are in providing text-to-image generators that produce trustworthy, accurate, and unbiased images when prompted to visualize different information topics or messages.

Based on theories about AI image generation, AI content moderation, AI generation errors, and AI biases, four main characteristics of AI imagery are theoretically discussed and measured in content analyses (e.g., Ahmad et al., 2026; Döring et al., 2025; Döring & Shevtsova, 2026; Mehrabi et al., 2022):

- **AI image over-blocking:** Blocking means the refusal of an AI image generator to create the requested image, returning a textual refusal instead. Over-blocking occurs when an AI image generator incorrectly refuses to generate an image in response to a legitimate prompt. Over-blocking can be an issue in many relevant information contexts where AI moderation systems do not adequately differentiate between harmful content (e.g., a sexualized image of a child) and harmless or helpful content (e.g., an image for a campaign against child marriage), and hence censor relevant visualizations, for example in fields such as sexual education, reproductive health, or the prevention of gendered violence (e.g., Vargas Penagos, 2024).
- **AI image flaws:** AI image generators can produce flawed images that do not accurately represent the requested content. These flaws fall into two categories: visual errors (e.g., distorted body parts, non-existent objects, or illogical spatial relationships) and – for images with textual elements such as diagrams – textual errors (e.g., factually incorrect or misspelled image labels). Both are examples of what researchers call AI hallucinations or confabulations – a well-documented problem where AI models produce content that may look convincing at first glance but is factually wrong (e.g., Amri & Hisan, 2023). This happens because AI generates what is statistically probable, not what is accurate. Users who share AI-generated images without fact-checking all visual and textual details therefore risk spreading false or misleading visuals.
- **AI image representational biases:** The systematic over- or underrepresentation of certain social groups in AI-generated images

depicting people is referred to as representational bias. It concerns whether the full diversity of humans is adequately represented. Because AI models are trained on data that reflects and entrenches existing societal biases, they are prone to reproducing and even amplifying them – for example, by generating predominantly White men in response to an occupational prompt such as the request to visualize “a doctor” (e.g., Hartley & Fisher, 2026). Representational biases are associated with representational harms, as inadequate representation can reinforce the marginalization of underrepresented groups.

- **AI image presentational biases:** Presentational bias refers to how different population groups are portrayed in AI-generated images, with a particular focus on visual stereotypes, such as ethnic, age, or gender stereotypes. For the same reason that AI models reproduce representational biases, they can also reproduce or amplify presentational biases. For example, in line with cultural gender stereotypes (Döring, 2022), AI-generated images depicting women tend to visualize them in domestic roles or as conventionally attractive, slim, young, smiling, and sexualized (e.g., Sun et al., 2024). Presentational biases are associated with presentational harms, as stereotyped presentations can reinforce discrimination against stereotyped groups.

AI ethics and AI regulations require AI companies to ensure AI fairness, meaning that AI decisions and AI-generated content should avoid biases. However, defining what constitutes an unbiased image is often challenging. For example, if AI mainly represents “a surgeon” as a conventionally attractive, able-bodied, White male doctor, this may be considered an instance of both representational and presentational biases (Ali et al., 2024; Cevik et al., 2024). Generating more diverse images of surgeons may be viewed as fairer because it better reflects the real diversity of surgeons worldwide, and may encourage a broader range of people to aspire to such careers, including for example non-White women. However, highly diverse AI-generated images can also be criticized as misleading because they may obscure existing structural inequalities. For instance, they may conceal the reality that academic positions in many Western countries are still disproportionately occupied by White, non-



disabled men. This raises the question of what exactly constitutes “fair AI representation and presentation”: Should AI systems accurately reflect current societal demographics and social roles, or should they increase the visibility of minority groups to promote inclusion and empowerment? AI companies implementing de-biasing filters to promote diversity (e.g., by avoiding representing professionals mainly as White men) have found that such interventions can lead to new mistakes. For example, users requesting images of German soldiers during World War II received images depicting Black or Asian women in German military uniforms, despite the historical inaccuracy of such representations (Robertson, 2024). Such outputs are not only misleading but may also be considered offensive when they inaccurately depict real historical events or groups. As a result, both the theoretical conceptualization of AI biases and related de-biasing measures, as well as their implementation in AI models, remain active areas of research. These considerations should be taken into account when analyzing and criticizing AI images for potential biases (Waters & Honenberger, 2025).

Beyond their AI-specific characteristics, AI-generated images are also often analyzed through the lens of Framing Theory (Rodriguez & Dimitrova, 2011; Thelander, 2018). Just as with other types of images, communication researchers may be interested in how certain topics – such as health conditions, treatment types, or hygiene products – are visually framed in AI-generated images. More specifically, this framing can be examined across different dimensions, such as the use of colors, symbols, and the denotative and connotative meanings of images (Döring & Shevtsova, 2026). Of particular interest is the overall valence of this framing – for example, whether AI-generated images frame menstruation as something natural and positive or as something negative and taboo. Similarly, AI-generated images on sexual health can be examined for sex-positive messages (Maes & Vandenbosch, 2024).

REFERENCES/COMBINATION WITH OTHER METHODS OF DATA COLLECTION

Content analyses of AI-generated images provide valuable insights into the visual characteristics and representations found in AI-generated imagery. The analyzed images may be collected from social media platforms, such as X/Twitter

(Asperti et al., 2025) or Instagram (e.g., Liu et al., 2026; Peng et al., 2025), or generated by researchers for specific topics and intended applications, such as news reporting (e.g., Elhosary, 2026), educational materials (e.g., Campbell et al., 2026), or health-related advertisements (e.g., Wang & Wang, 2025). Such content analysis studies may be qualitative (e.g., Wang & Wang, 2025) or quantitative and may rely on either manual coding (e.g., Reed et al., 2025) or computational analysis techniques (e.g., Sun et al., 2024). Some studies combine qualitative and quantitative approaches to capture both large-scale patterns and more nuanced visual characteristics (e.g., Högemann et al., 2025; Peng et al., 2025). Content analyses of AI-generated images can also be combined with other methods to investigate image creators’ experiences with AI image production and image recipients’ perceptions, use, and perceived effects of AI-generated images, for example, students’ and teachers’ evaluations of the use of AI-generated images in schools. These complementary methods include interviews (e.g., Mahdavi Goloujeh et al., 2024), group discussions (e.g., Reed et al., 2023), and surveys (e.g., Bastý et al., 2025). Some studies have additionally used experimental designs to compare participants’ perceptions of AI-generated versus human-generated images (e.g., Döring & Mohseni, 2025; Marini et al., 2024) and to examine the effects of using AI-generated images in different application domains, such as education (e.g., Kang et al., 2025; Romanik & Wegner, 2025).

EXAMPLE STUDIES FOR MANUAL QUANTITATIVE CONTENT ANALYSES

Several content analysis studies have examined informational AI-generated images; in the field of health communication alone, twenty manual quantitative content analyses were identified (Döring & Shevtsova, 2026). These studies can be grouped into three main thematic areas: AI-generated images of (1) health professionals, (2) patients, and (3) populations with particular health relevance. For example, Reed et al. (2025) analyzed N = 288 AI-generated images of nurses produced by three image generators (Adobe Firefly, ImageFX, and Midjourney). Similarly, Chou et al. (2026) analyzed N = 320 AI-generated images depicting cancer patients created with DALL-E and Stable Diffusion. Sobey (2025) examined the representation of different body sizes through an

analysis of N = 649 images generated by nine AI image models.

The current example study (Döring & Shevtsova, 2026) applies manual quantitative content analysis to examine the characteristics of AI-generated images related to sexual and reproductive health and rights (SRHR) across 15 AI image generators, with a particular focus on images depicting menstruation for an educational brochure. The study is grounded in concepts drawn from theories about AI image generation, AI content moderation, AI generation errors, and AI biases – mainly representational and presentational biases and the associated notions of representational and presentational harms. The analysis is based on a sample of N = 600 menstruation-related AI-generated outputs. The codebook was developed deductively based on the research questions and the existing literature on the characteristics of

AI-generated health images and was subsequently refined inductively using example material. The codebook was evaluated for its three key quality criteria: objectivity, validity, and reliability (Döring, 2023): To ensure objectivity, all concepts are explained in detail with coding instructions and visual examples, and the full codebook is publicly available (<https://osf.io/h5pn3/>). To ensure validity, all categories in the codebook are grounded in the theoretical and empirical literature with respective references. To ensure reliability, all categories were tested for intercoder reliability between two independent trained coders using a subsample of 60 images. A codebook excerpt with concept definitions, operationalizations, reliability coefficients, and references to content analyses measuring similar variables can be found in Table 1.

Table 1. Codebook to Identify Characteristics of AI-generated Information Images

Characteristics of AI-generated images	Operationalizations (Döring & Shevtsova, 2026)	Reliability (Gwet's AC1 coefficient)	References to further content analyses measuring similar variables
(1) AI Image Over-Blocking: Incorrect refusal to generate images for legitimate prompts	Prompt outcome (only for studies with researcher-generated images using legitimate prompts): Over-Blocking: code 0 = AI output is an image / code 1 = AI output is not an image but a blocking message	1.00	Over-Blocking: Rando et al. (2022); Westberg & Kvåle (2025, pp. 20-21)
(2) AI Image Flaws: Visual and textual manifestations of AI hallucinations / confabulations	Occurrence of visual and/or textual flaws: - Visual errors (for all images): code 0 = no / code 1 = yes (image contains at least one visual error) - Textual errors (only for images with textual elements): code 0 = no / code 1 = yes (image with textual elements contains at least one textual error)	.86 .96	Visual errors: Campbell et al. (2026) Textual errors: Temsah et al. (2024)

**(3) AI Image Representational Bias:**

Over- and/or under-representation of population groups

Representation of different population groups (only for images that depict at least one person; representation can be measured either for the selected focal person in the image or for all persons in the image by numbering them and duplicating the variables in the form: Gender_Person1; Gender_Person2 etc.):

- Gender: code 1 = female-appearing person / code 2 = male-appearing person / code 3 = non-binary or gender-diverse-appearing person / code 99 = gender not codable
- Age: code 1 = child / code 2 = adolescent / code 3 = adult / code 4 = older adult / code 99 = age not codable
- Skin Tone: code 1 = lighter skin tone / code 2 = medium skin tone / code 3 = darker skin tone / code 99 = skin tone not codable
- Body Type: code 1 = slim body type / code 2 = mid-size body type / code 3 = plus-size or larger body type / code 99 = body type not codable

.88
.88
.87
.90

Gender: Campbell et al. (2026)

Age: Gisselbaek et al. (2024)

Skin tone: Elhosary (2026)

Body type: Elhosary (2026)

(4) AI Image Pre-sentational Bias:

Stereotypical portrayals of people

Stereotyped presentation of women (only for images that depict at least one female-appearing person; stereotyped presentation can be measured either for the selected focal woman in the image or for all women in the image by numbering them and duplicating the variables in the form: Sexualized-Appearance_Woman1; Sexualized-Appearance_Woman2 etc.)

- Sexualized Appearance: code 0 = no / code 1 = yes (depicted woman appears sexualized)
- Smiling: code 0 = no / code 1 = yes (depicted woman is smiling)
- Long Hair: code 0 = no / code 1 = yes (depicted woman has long hair)
- Dress/Skirt: code 0 = no / code 1 = yes (depicted woman wears a dress or skirt)

.93
.86
.96
1.00

Sexualized appearance: Reed et al. (2025)

Smiling: Reed et al. (2025)

Long hair: Westberg & Kvåle (2025)

Dress/skirt: Westberg & Kvåle (2025)

Depending on the research field and research questions, the coding variables can be further expanded. For example, visual flaws may be differentiated into more specific categories, such as anatomical errors or non-existent objects, while textual flaws may be further classified into sub-categories such as spelling and grammar errors or informational inaccuracies (Döring & Shevtsova, 2026). Likewise, additional demographic and personal characteristics can be included to examine representational biases in the representation of people, such as religion, socioeconomic

status, sexual orientation, or disability status (Alon et al., 2026). Moreover, presentational biases are not limited to gender stereotyping but can also be examined in terms of ethnic, national, age, or disability stereotypes. Importantly, stereotypical portrayals do not only concern the presentation of persons but also their environments. For example, AI-generated science classrooms have been found to be depicted mainly as well-equipped laboratories, thereby evoking the image of elite educational environments (Cooper & Tang, 2024).

REFERENCES

- Ahmad, A., Vallès, Y., & Idaghdour, Y. (2026). Bias in AI systems: Integrating formal and socio-technical approaches. *Frontiers in Big Data*, 8, 1686452. <https://doi.org/10.3389/fdata.2025.1686452>
- Ali, R., Tang, O. Y., Connolly, I. D., Abdulrazeq, H. F., Mirza, F. N., Lim, R. K., Johnston, B. R., Groff, M. W., Williamson, T., Svokos, K., Libby, T. J., Shin, J. H., Gokaslan, Z. L., Dobberstein, C. E., Zou, J., & Asaad, W. F. (2024). Demographic Representation in 3 Leading Artificial Intelligence Text-to-Image Generators. *JAMA Surgery*, 159(1), 87. <https://doi.org/10.1001/jamasurg.2023.5695>
- Alon, L., Hadar Shoval, D., & Levkovich, I. (2026). Bias and representation in AI generated text-to-image in education: A systematic review. *Computers and Education: Artificial Intelligence*, 10, 100587. <https://doi.org/10.1016/j.caeai.2026.100587>
- Amri, M., & Hisan, U. (2023). Incorporating AI Tools into Medical Education: Harnessing the Benefits of ChatGPT and Dall-E. *Journal of Novel Engineering Science and Technology*, 2(02), 34–39. <https://doi.org/10.56741/jnest.v2i02.315>
- Asperti, A., George, F., Marras, T., Stricescu, R. C., & Zanotti, F. (2025). A Critical Assessment of Modern Generative Models' Ability to Replicate Artistic Styles. *Big Data and Cognitive Computing*, 9(9), 231. <https://doi.org/10.3390/bdcc9090231>
- Basty, R., Kropczynski, J., & Halse, S. (2025). Exploring Higher Education Faculty Insights on Generative AI in Creative Courses. *Journal of Information Technology Education: Research*, 24, 018. <https://doi.org/10.28945/5546>
- Buzzaccarini, G., Degliuomini, R. S., Borin, M., Fidanza, A., Salmeri, N., Schiraldi, L., Di Summa, P. G., Vercesi, F., Vanni, V. S., Candiani, M., & Pagliardini, L. (2024). The Promise and Pitfalls of AI-Generated Anatomical Images: Evaluating Midjourney for Aesthetic Surgery Applications. *Aesthetic Plastic Surgery*, 48(9), 1874–1883. <https://doi.org/10.1007/s00266-023-03826-w>
- Campbell, L. O., Lambie, G. W., Hayes, B. G., & Hartshorne, R. (2026). Generative Artificial Intelligence Images for Instruction: A Content Analysis. *TechTrends*, 70(1), 80–92. <https://doi.org/10.1007/s11528-025-01129-2>
- Cevik, J., Lim, B., Seth, I., Sofiadellis, F., Ross, R. J., Cuomo, R., & Rozen, W. M. (2024). Assessment of the bias of artificial intelligence generated images and large language models on their depiction of a surgeon. *ANZ Journal of Surgery*, 94(3), 287–294. <https://doi.org/10.1111/ans.18792>
- Chou, W.-Y. S., Gaysynsky, A., Senft-Everson, N., Muro, A., Schrader, K., & Iles, I. (2026). Portrayal of Cancer Patients in the Era of AI: A Content Analysis of Images Produced by Generative AI Tools. *Health Communication*, 41(5), 888–898. <https://doi.org/10.1080/10410236.2025.2537807>
- Cooper, G., & Tang, K.-S. (2024). Pixels and Pedagogy: Examining Science Education Imagery by Generative Artificial Intelligence. *Journal of Science Education and Technology*, 33(4), 556–568. <https://doi.org/10.1007/s10956-024-10104-0>
- Döring, N. (2022). Visual Gender Stereotypes (Advertisement, Social Media). DOCA – Database of Variables for Content Analysis, 1(5). <https://doi.org/10.34778/5i>

- Döring, N. (2023). *Forschungsmethoden und Evaluation in den Sozial- und Humanwissenschaften* [Research Methods and Evaluation in the Social and Human Sciences] (6th ed.). Springer.
- Döring, N., & Mohseni, M. R. (2025). Anti-AI Bias Toward Couple Images and Couple Counseling: Findings from Two Experiments. *Archives of Sexual Behavior*. <https://doi.org/10.1007/s10508-025-03318-9>
- Döring, N., & Shevtsova, A. (2026). Visual Sexual Health Communication Assisted by AI: The Case of Synthetic Menstruation Images. *Archives of Sexual Behavior*. <https://doi.org/10.1007/s10508-026-03518-x>
- Döring, N., Le, T. D., Vowels, L. M., Vowels, M. J., & Marcantonio, T. L. (2025). The Impact of Artificial Intelligence on Human Sexuality: A Five-Year Literature Review 2020–2024. *Current Sexual Health Reports*, 17(1), 4. <https://doi.org/10.1007/s11930-024-00397-y>
- Elhosary, M. (2026). Pixels of Prejudice: Decoding Embedded Biases in AI-Generated News Imagery and Their Implications for Visual Journalism—Toward an Algorithmic-Mediated Visual Framing. *Journalism & Mass Communication Quarterly*, 103(2), 430–461. <https://doi.org/10.1177/10776990251360690>
- Gisselbaek, M., Minsart, L., Köseleli, E., Suppan, M., Meco, B. C., Seidel, L., Albert, A., Barreto Chang, O. L., Saxena, S., & Berger-Estilita, J. (2024). Beyond the stereotypes: Artificial Intelligence image generation and diversity in anesthesiology. *Frontiers in Artificial Intelligence*, 7, 1462819. <https://doi.org/10.3389/frai.2024.1462819>
- Gunasekara, L., El-Haber, N., Nagpal, S., Moraliyage, H., Issadeen, Z., Manic, M., & De Silva, D. (2025). A Systematic Review of Responsible Artificial Intelligence Principles and Practice. *Applied System Innovation*, 8(4), 97. <https://doi.org/10.3390/asi8040097>
- Hartley, A., & Fisher, J. (2026). Seeing is believing? Exploring gender bias in artificial intelligence imagery of specialty doctors. *The Clinical Teacher*, 23, e70297. <https://doi.org/10.1111/tct.70297>
- Högemann, M., Betke, J., & Thomas, O. (2025). What you see is not what you get anymore: A mixed-methods approach on human perception of AI-generated images. *Frontiers in Artificial Intelligence*, 8, 1707336. <https://doi.org/10.3389/frai.2025.1707336>
- Kang, W., Lee, H.-S., & Hong, J.-H. (2025). AI-assisted design communication in ceramic-crafts education: Investigating image generation tools for concept representation and pedagogical practice. *Computers & Education*, 239, 105419. <https://doi.org/10.1016/j.compedu.2025.105419>
- Lapointe, V. A., Dubé, S., Rukhlyadyev, S., Kesai, T., & Lafortune, D. (2025). The Present and Future of Adult Entertainment: A Content Analysis of AI-Generated Pornography Websites. *Archives of Sexual Behavior*. <https://doi.org/10.1007/s10508-025-03099-1>
- Liu, X., Lu, Y., Peng, Q., Qian, S., Peng, Y., & Shen, C. (2026). Seeing the Surreal: Mapping Surrealism in Photorealistic AI-Generated Images Using Large Language Models. *Computational Communication Research*, 8(2). <https://doi.org/10.5117/CCR2026.2.6.LIU>
- Maes, C., & Vandenbosch, L. (2024). Positive Sexuality. *DOCA – Database of Variables for Content Analysis*, 1(3). <https://doi.org/10.34778/3k>
- Mahdavi Goloujeh, A., Sullivan, A., & Magerko, B. (2024). Is It AI or Is It Me? Understanding Users' Prompt Journey with Text-to-Image Generative AI Tools. *Proceedings of the CHI Conference on Human Factors in Computing Systems*, 1–13. <https://doi.org/10.1145/3613904.3642861>
- Marini, M., Ansani, A., Demichelis, A., Mancini, G., Paglieri, F., & Viola, M. (2024). Real is the new sexy: The influence of perceived realness on self-reported arousal to sexual visual stimuli. *Cognition and Emotion*, 38(3), 348–360. <https://doi.org/10.1080/02699931.2023.2296581>
- Mehrabi, N., Morstatter, F., Saxena, N., Lerman, K., & Galstyan, A. (2022). A Survey on Bias and Fairness in Machine Learning. *ACM Computing Surveys*, 54(6), 1–35. <https://doi.org/10.1145/3457607>
- Peng, Q., Lu, Y., Peng, Y., Qian, S., Liu, X., & Shen, C. (2025). Crafting synthetic realities: Examining visual realism and misinformation potential of photorealistic AI-generated images. In *CHI EA '25: Proceedings of the Extended Abstracts of the CHI Conference on Human Factors in Computing Systems* (Article 156, pp. 1–12). ACM. <https://doi.org/10.1145/3706599.3719834>

- Rando, J., Paleka, D., Lindner, D., Heim, L., & Tramèr, F. (2022). Red-teaming the Stable Diffusion safety filter. Workshop on Machine Learning Safety, 36th Conference on Neural Information Processing Systems (NeurIPS 2022). <https://doi.org/10.48550/arXiv.2210.04610>
- Reed, J., Alterio, B., Coblenz, H., O'Lear, T., & Metz, T. (2023). AI Image-Generation as a Teaching Strategy in Nursing Education. *Journal of Interactive Learning Research*, 34(2), 369–399. <https://doi.org/10.70725/729255gdiorw>
- Reed, J., Dodson, T., Petrincec, A., Tennant, D., Chmelik, J., & Cripple, S. (2025). Artificial Intelligence and Images Portraying Nurses Through the Decades. *OJIN: The Online Journal of Issues in Nursing*, 30(2). <https://doi.org/10.3912/OJIN.Vol30No02Man05>
- Robertson, A. (2024, February 21). Google apologizes for 'missing the mark' after Gemini generated racially diverse Nazis. *The Verge*. <https://www.theverge.com/2024/2/21/24079371/google-ai-gemini-generativeinaccurate-historical>
- Rodriguez, L., & Dimitrova, D. V. (2011). The levels of visual framing. *Journal of Visual Literacy*, 30(1), 48–65. <https://doi.org/10.1080/23796529.2011.11674684>
- Romanik, M., & Wegner, C. (2025). Effects of videos with AI-generated images to provide students with authentic and diverse insights into STEM occupations. *American Journal of STEM Education*, 5, 15–34. <https://doi.org/10.32674/edckb827>
- Sanseviero, O., Cuenca, P., Passos, A., & Whitaker, J. (2024). Hands-on generative AI with transformers and diffusion models. O'Reilly Media.
- Sobey, A. (2025). The thinness of GenAI: Body size in relation to the construction of the normate through GenAI image models. *AI and Ethics*, 5(4), 4181–4196. <https://doi.org/10.1007/s43681-025-00684-x>
- Sun, L., Wei, M., Sun, Y., Suh, Y. J., Shen, L., & Yang, S. (2024). Smiling women pitching down: Auditing representational and presentational gender biases in image-generative AI. *Journal of Computer-Mediated Communication*, 29(1), zmad045. <https://doi.org/10.1093/jcmc/zmad045>
- Temsah, M.-H., Alhuzaimi, A. N., Almansour, M., Aljamaan, F., Alhasan, K., Batarfi, M. A., Altamimi, I., Alharbi, A., Alsuhailbani, A. A., Alwakeel, L., Alzahrani, A. A., Alsulaim, K. B., Jamal, A., Khayat, A., Alghamdi, M. H., Halwani, R., Khan, M. K., Al-Eyadhy, A., & Nazer, R. (2024). Art or Artifact: Evaluating the Accuracy, Appeal, and Educational Value of AI-Generated Imagery in DALL·E 3 for Illustrating Congenital Heart Diseases. *Journal of Medical Systems*, 48(1), 54. <https://doi.org/10.1007/s10916-024-02072-0>
- Thelander, Å. (2018). Visual Communication. In R. L. Heath & W. Johansen, *The International Encyclopedia of Strategic Communication* (1st ed., pp. 1–13). Wiley. <https://doi.org/10.1002/9781119010722.iesc0198>
- Vargas Penagos, E. (2024). ChatGPT, can you solve the content moderation dilemma? *International Journal of Law and Information Technology*, 32, eaae028. <https://doi.org/10.1093/ijlit/eaae028>
- Wang, T., & Wang, R. (2025). What AI knows about breast cancer: An exploratory visual discourse analysis on AI-generated breast cancer advertisement images. *Social Semiotics*, 1–37. <https://doi.org/10.1080/10350330.2025.2519471>
- Waters, G., & Honenberger, P. (2025). AI biases as asymmetries: A review to guide practice. *Frontiers in Big Data*, 8, 1532397. <https://doi.org/10.3389/fdata.2025.1532397>
- Westberg, G., & Kvåle, G. (2025). The generic uniqueness of AI imagery: A critical approach to Dall-E as semiotic technology. *Discourse & Society*, 36(4), 575–598. <https://doi.org/10.1177/09579265241274788>

FUNDING

This article was prepared with support from the network 'Interdisciplinary Pornography Research', funded by the German Research Foundation (Deutsche Forschungsgemeinschaft, DFG; project number 568648320). The first author, Nicola Döring, is a member of the network.