



Carlos Monteiro

University of Coimbra
carlos_kh2fm@hotmail.com

Alexandra Kuzmina

Erasmus University Rotterdam
Ghent University
alexkuzmina.work@gmail.com

THE RHETORIC OF PROMPTS: MULTIMODALITY, INTENT, AND ETHICS IN AI MEANING-MAKING

<https://doi.org/10.17234/9789533793337.20>

Abstract

The rise of generative AI poses new challenges to public discourse, particularly through synthetic media whose origin, authorship, and intent are often ambiguous. These challenges extend beyond misinformation to the underlying structures through which meaning is constructed, framed, and ethically justified. Addressing them requires a conceptual framework that treats AI-generated content not as neutral or purely technical output, but as rhetorical artefacts embedded within broader ethical and communicative systems. This article proposes a metatheoretical approach that draws on rhetoric and multimodality to reframe AI-generated content – specifically AI-generated imagery – as sites of ethical activity and moral accountability. Central to our argument is the notion that rhetorical intent can still be identified in the human agents who design, prompt, and disseminate these outputs. Prompts function as rhetorical acts that encode thematic, stylistic, and ideological directives, translated across semiotic modes into persuasive visual forms. Building on foundational theories in rhetoric (Perelman and Olbrechts-Tyteca 1969), multimodality (Kress and Leeuwen 2001; Poulsen 2015), and visual argumentation (Kjeldsen 2015; Žagar 2020), we conceptualize intentionality as a distributed process enacted through user-system interaction, thus enabling a richer understanding of how values, identities, and ideologies are encoded, circulated, and interpreted in AI-mediated communication, especially within hybrid public discourse.

Keywords: AI ethics, argumentation, generative content, language models, multimodality, rhetoric



1. Introduction

A growing body of research warns that AI-generated content may be exploited to manipulate perception, manufacture consensus, or distort political and social realities (cf. Fisher, Howard, and Kira 2024). But while scholarship on algorithmic bias has advanced considerably (cf. Lendvai and Gosztönyi 2025; Hadley, Blatecky, and Comfort 2025), critical analysis of AI-generated media – particularly visual content – remains comparatively underdeveloped. This gap can be attributed, in part, to the inadequacy of traditional frameworks of authorship, intention, and accountability when applied to hybrid outputs produced through a combination of user prompts and algorithmic inference.

In conventional media, authorial intent is typically reconstructible through the conscious, communicative acts embedded in the creative process – what Jürgen Habermas would describe as speech acts oriented toward mutual understanding (Habermas 1984, 110–29). Such works reflect the author's intention to convey meaning, subject to normative claims of truthfulness, sincerity, and appropriateness. By contrast, generative AI systems produce outputs through probabilistic pattern recognition across large-scale datasets, modulated by textual prompts that may or may not fully reflect the user's ethical, rhetorical, or intentional aims. This entanglement of human agency and machinic computation raises pressing questions: Where does intent reside in AI-generated content? Who bears responsibility for the meanings or persuasive effects such content may produce? Can we establish normative frameworks capable of addressing ethical accountability in the production and circulation of AI-generated media?

These questions become particularly salient in the realm of visual content. Compared to linguistic discourse, images often carry a greater degree of ambiguity regarding their meaning (Barthes 1977, 32) – yet it is precisely this ambiguity that grants them rhetorical power, as one of the central strengths of visual media lies in its capacity to transcend linguistic and cultural boundaries: a single image can evoke universal themes – defiance, victimhood, martyrdom, power – making it a potent instrument in the rhetoric of global movements (McClanahan 2023, 12–50). This visual “universality” is especially significant in today's fragmented online environment, where audiences from diverse sociocultural backgrounds may nonetheless converge around shared grievances or ideological affinities.

AI-generated imagery introduces unique analytical challenges, particularly because its creation diverges significantly from traditional image-making. Rather than being drawn, painted, or assembled by hand, these visuals are produced by generative models that *convert textual prompts into visual output*. Prompts function as the interface through which the system interprets and recombines vast corpora of pre-existing imagery. As a result, the final image does not arise from a user's individual creative *poiesis*, but from large-scale pattern recognition. Yet users still play an active curatorial role: they evaluate the system's output for coherence and alignment with their intended message, often refining prompts through multiple iterations. What ultimately circulates online is typically the result of this iterative, user-guided process.

In the context of AI-mediated communication, the question of intention is not merely speculative – it bears directly on debates around liability (in law and criminology), authorship (in media studies), and agency (in anthropology). As Ted Chiang provocatively suggests, generative AI “reduces the amount of intention in the world” by automating decisions that would otherwise require deliberative human agency (Chiang 2024). This challenges longstanding philosophical accounts of intention as a consciously enacted mental state expressed through action (cf. Aristotle 2000, 37–59; Kant 1997, 1–16; Anscombe 1963, 20–50; Davidson 1978, 41–50). To engage this tension, this article proposes a metatheoretical framework that brings rhetoric and multimodality



into dialogue, enabling a nuanced identification and interpretation of the intentionality embedded in AI-generated content. Central to this framework is the understanding of prompts as *deliberate, intentional acts* and the resulting outputs as *rhetorically charged expressions with ethical implications*. In doing so, this approach lays the groundwork for theorizing ethical accountability that encompasses and is attentive to the intricate moral issues involved in the creation, dissemination, and reception of AI-generated media.

2. Intent, Rhetoric, and Meaning

Our philosophical starting point is Chaïm Perelman and Lucie Olbrechts-Tyteca's *The New Rhetoric: A Treatise on Argumentation*, a work that had a foundational impact in reconceptualizing rhetoric beyond classical Aristotelian models. According to Perelman and Olbrechts-Tyteca, the study of rhetoric, at its core, is the study of how meaning is constructed, contested, and communicated with the intent to persuade. Rhetoric, in this view, is not merely concerned with the transmission of information but with *framing*: shaping how ideas are interpreted, which values are foregrounded, and what inferences an audience is likely to draw. Persuasion, thus, is not reducible to speech alone; it is a relational act grounded in the implicit negotiation of meaning between speaker and audience (Perelman and Olbrechts-Tyteca 1969, 14–31, 65–70).

This conception of rhetoric aligns naturally with the principles of multimodality, which reject the primacy of linguistic form and instead emphasize that communication is constituted by the interplay of multiple semiotic modes – visual, textual, gestural, spatial, and auditory (Flewitt, Price, and Korkiakangas 2023; Kress and Leeuwen, 2001, 1–3, 20–24; Stöckl and Tsenoris 2024, 167–72; Forceville 2009, 382–83). From a multimodal perspective, *meaning* is emergent and situated: it arises not from isolated symbols but from the structured coordination of modes within specific contexts (Poulsen 2015). Crucially, multimodal theorists argue that visual or spatial cues often do the rhetorical “work” that language alone cannot accomplish, particularly in highly compressed or emotionally charged content such as propaganda, advertising, or meme culture (Bülow and Johann 2023; Kjeldsen 2016, 266–72). This emphasis on situated meaning and audience engagement is closely aligned with Perelman and Olbrechts-Tyteca's “new rhetoric”, which similarly foregrounds the contextual and value-laden nature of persuasion. Both frameworks emphasize that rhetorical effectiveness (i.e. persuasion) depends on how meaning is shaped, negotiated, and rendered persuasive through the resources available in a given communicative situation (Pflaeging and Stöckl 2021; Johnstone 1990; Burke 1969, 49–64).

In the context of generative AI, a dual framework that draws on *both* rhetoric and multimodality is especially pertinent. Models such as GPT and Stable Diffusion are not merely computational tools; their outputs are generated from massive training datasets that reflect specific cultural, political, and historical conditions – none of which are ideologically neutral (Bender et al. 2021, 610–15; Crawford 2021, 4–14). The curation of training data – what is included, excluded, and how it is labeled – constitutes a form of rhetorical selection, shaping the range of meanings these systems are capable of producing. In short, they are discursive systems that organize knowledge, encode ideologies, and shape public perception.

But generative models do not simply reflect the world; *they participate in constructing it*. Although they lack intent, they nonetheless function as agents of framing, shaping how audiences encounter and interpret meaning through probabilistic language generation. In doing so, they contribute to the *axiologization* of discourse – assigning value to ideas, legitimizing certain narratives,



and normalizing particular associations. This presents a paradox: rhetorical influence is seemingly exercised through systems without agency, intention, or accountability. To address this paradox, we must reorient our analytical focus: from system to user, from output to prompting, and from authorship to *delegated* intention.

The distribution of persuasive force across this human-machine assemblage demands a rethinking of how we theorize authorship, responsibility, and communication ethics. If we accept that language – and by extension, generated language and imagery – is a primary tool for worldmaking (whether individually, socially, or politically), then the rhetorical stakes of AI systems become unmistakable. Those who engineer, fine-tune, and deploy these systems exert disproportionate influence over the prescriptive dimensions of meaning: what is persuasive, what is believable, and even what is *thinkable*. In this sense, AI systems function not merely as tools of automation, but as infrastructures of epistemic power. Analyzing AI systems through a rhetorical and multimodal lens allows us to foreground this power and critically examine its effects – not only in terms of bias or error, but in terms of the narratives and values they normalize.

3. The Problem of Intentionality in AI-generated Content

When analyzed through a rhetorical/multimodal lens, the first thing to note is that AI-generated content diverges fundamentally from classical rhetorical models, which typically rely on a triadic structure: *speaker*, *audience*, and *message* (Bitzer 1968, 1–6; Fidalgo 2009, 232). These models presuppose an intentional speaker – someone with beliefs, goals, and a communicative purpose – who strategically crafts discourse to influence a target audience (Aristotle 2012, 5–19). In contrast, generative AI lacks a unified speaker; the message it produces is the result of stochastic pattern-matching across vast training corpora. This absence of communicative intent raises a fundamental question: *Is rhetorical meaning possible in the absence of a rhetor?*

Rhetorically speaking, communication presupposes at least minimal intentionality, even if that intent is implicit, contested, or fragmented (Charland 1987, 137–40). However, language models and image generators possess no beliefs, intentions, or awareness of an audience. They neither respond to feedback nor “*mean*” anything in a semantic sense (Bender et al. 2021, 611, 615–18; Dennett 1988, 180–98). Their outputs simulate coherence and relevance via probabilistic approximations of prior linguistic patterns. Such systems are “*stochastic parrots*” mimicking the surface form of language without accessing its semantic depth. The result is a kind of *pseudo-discourse* – language that appears meaningful but originates from systems incapable of communicative intent (Bender et al. 2021, 616–19).

However, if we understand prompt engineering as a form of performative utterance, it becomes legible as what Perelman and Olbrechts-Tyteca would describe as a rhetorical act (Perelman and Olbrechts-Tyteca 1969, 14–21). In this view, the prompt functions not merely as a descriptive instruction but as an illocutionary act, intended to elicit a visual or textual response aligned with the user’s communicative goals. Thus, prompt-writing can be understood as a rhetorical practice – one that encodes thematic, stylistic, and ideological directives (Perelman and Olbrechts-Tyteca 1969, 115–25, 142–53). When a user instructs an AI to generate an image of “a politician as a monster” or “a holy warrior in flames”, they are not merely issuing a command; they are designing a communicative event, shaping both the argumentative and affective contours of the output. The prompt projects an intention that will later be judged based on the perceived coherence between the input and the result. This means that AI-generated materials are neither ethically neutral nor



rhetorically inert. They emerge from *distributed intention*: shaped by human judgment, filtered through machine learning architectures, and, ultimately, interpreted by audiences through socially contingent lenses.

This model of communication – fragmented, mediated, and iterative – requires a reconceptualization of rhetorical agency. Rather than locating agency in a single, authorial figure, we must trace it across a chain of actions: prompt composition, model inference, output selection, and public dissemination. Only by following this chain can we begin to theorize accountability in communicative environments where authorship is partial, collaborative, and technically opaque.

4. The Rhetoric of AI-generated Images

However, the transformation of a textual prompt into a visual output introduces a cross-modal rhetorical dynamic. In this process, linguistic discourse becomes the interface through which users activate latent visual forms encoded within the model's training data. But this is not merely a case of description leading to depiction; it is an act of *semiotic translation*. The resulting image bears the imprint of *both* user intention and algorithmic bias, forming a hybrid communicative artifact. Even when the prompt remains invisible to downstream audiences, its influence is embedded in the structure, tone, and ideological resonance of the generated image.

The rhetorical function of visual media has long been underestimated within traditions that privilege verbal discourse (Cardoso e Cunha 2009, 223). Yet as digital environments become increasingly image-dominant, visual forms have emerged not as supplements to argumentation, but, some would argue, as its primary vehicles (Delicath and DeLuca 2003, 315–25; Buhel 2016, 3–12). This is especially evident in the case of AI-generated images, which are constructed through textual prompts but often encountered without any accompanying linguistic context (Alikhani, Khalid, and Stone 2023).

This shift compels a reexamination of foundational questions about visual meaning. Do images inherently carry meaning, or do they become legible only through linguistic or cultural codes? Roland Barthes' use of the distinction between *denotation* (what is shown) and *connotation* (what is meant) offers a useful starting point (Barthes 1977, 42–46) – but AI-generated imagery complicates this binary. These images are produced through textual prompts, yet the originating text is often invisible to the viewer (Trinh et al. 2024). As a result, the connotative dimension of the image both *precedes* and *exceeds* it, being shaped by user intention, platform affordances, and algorithmic patterning. To fully grasp the persuasive force of such images, we must move beyond the question “*What does this image depict?*” and instead ask: “*What does this image do – and how does it do it?*”

The rhetorical power of an image resides primarily in its connotative function – its capacity to suggest, evoke, or insinuate rather than explicitly state (Barthes 1977, 46–49). AI-generated images excel within this connotative realm. They evoke ideological positions, frame social realities, and reinforce cultural myths – all without overt declaration. Their strength lies in suggestive ambiguity, inviting audiences to “complete” the argument themselves. This makes them especially potent in persuasive contexts such as political propaganda, ideological memes, symbolic martyrdom, or deepfakes. Such images often deploy shared visual grammars – halo effects, military aesthetics, and ethno-nationalist iconography – to generate emotional resonance while simultaneously disavowing authorial intent.

In traditional media, textual anchoring – through captions, headlines, or surrounding discourse – plays a crucial role in guiding interpretation (Barthes 1977, 37–41; Liu 2020). With AI-generated



imagery, however, this anchoring is displaced: the textual prompt that shapes the image is typically *invisible* to its audience, yet its influence is inscribed within the composition. A prompt such as “a Zionist leader as a fascist dictator” or “a martyr ascending to heaven” does not appear alongside the image but fundamentally structures the image’s affective and symbolic cues. Viewers are thus invited to infer intent through composition, symbolism, aesthetic framing, and cultural resonance. In this sense, AI-generated images are *pre-anchored* – not by visible text, but by invisible prompting.

These images represent a complex assemblage of human intent, machinic patterning, and audience interpretation. Their rhetorical power does not derive from fidelity to truth, but from their ability to mobilize symbols, shape inferences, and invoke shared cultural frames. To understand these images as persuasive acts requires more than visual literacy – it demands a theory of meaning that accounts for prompt authorship, algorithmic constraints, and multimodal inference.

4. Multimodal Argumentation and the “Visual Argumentation” Behind Meaning-making

Contemporary rhetorical and argumentation theories have increasingly emphasized multimodal persuasion, recognizing that arguments are not confined to verbal syntax but often emerge through the interaction of text, image, layout, and other semiotic modes (Kjeldsen 2015, 115–21). Visual elements can perform distinct argumentative functions – such as analogy, causality, or moral appeal. For instance, juxtaposing an AI-generated image of a crumbling cityscape with a smiling politician can visually imply that the politician is responsible. Crucially, these visual arguments do not rely on accompanying text for persuasiveness; rather, they depend on the audience’s inferential labour, cultural background, and visual literacy.

Multimodal argumentation is not only epistemic but also affective and symbolic. It goes beyond expressing opinions to actively constructing shared realities. (Mehu 2015) AI-generated imagery often operates rhetorically through visual enthymemes – implicit arguments where key premises are unstated but assumed by viewers (Kjeldsen 2015, 117–18). A visual depiction of a burning flag or broken bridge may not explicitly outline causality, but it strongly suggests themes of guilt, failure, or betrayal. Such images invite viewers to “fill in” the missing meaning, thus engaging them directly in the act of persuasion. In this sense, AI-generated visuals do not simply illustrate arguments – they *are* arguments, albeit fragmented and inherently multimodal (Gonçalves-Segundo et al. 2021, 722–27).

This dynamic is precisely why multimodality offers rhetoric and argumentation theory a crucial and refreshing framework for understanding how meaning is not simply stated but assembled through the interaction of diverse semiotic modes, each contributing uniquely to the communicative effect. Multimodality highlights this relational process by treating images, text, layout, and spatial organization as co-constitutive elements of meaning-making. Unlike structuralist models that view meaning as static and hierarchical, multimodal theory acknowledges communication as inherently relational and inferential (Flewitt, Price, and Korkiakangas 2018). Audiences do not decode messages mechanically; instead, they actively interpret by drawing inferences across modes. Visual arguments function not as mere illustrations but as inferential structures inviting the audience to “connect the dots” – to draw conclusions based on resemblance, contrast, spatial framing, and symbolic association (Pflaeging and Stöckl 2021). This inferential work is indispensable for understanding how persuasive meaning is constructed, especially in AI-generated visuals where linguistic anchoring is often absent.



From this perspective, AI-generated images are not merely outputs – they are argumentative acts. They function according to distinct patterns of visual reasoning: analogical, indexical, or iconic (Kuzmina 2024, 519–20; Chung et al. 2024; Noord and Garcia 2025). These arguments do not depend on deductive syllogisms but rather on a “visual argumentation” (Kjeldsen 2015, 115–16). For example, a composite image of a politician superimposed onto a battlefield invites viewers to infer associations with war, culpability, or betrayal – without any explicit statement. The rhetorical power derives from what is implied rather than overtly declared. In such cases, multimodality facilitates a form of enthymematic reasoning where audiences “fill in” unstated premises based on visual cues and shared cultural narratives (cf. Žagar 2020).

The generative process adds another layer of complexity. Prompts act as rhetorical blueprints, encoding user preferences for composition, affect, and ideology. These preferences are actualized through what Nataliia Laba terms “affordance activation” (Laba 2024, 5): users employ stylistic modifiers – such as “dramatic lighting,” “hyperreal,” or “military aesthetic” – to shape the aesthetic and emotional tone of the output. While the AI model generates images based on statistical patterning, the user exerts rhetorical control through prompt engineering. This hybrid interaction produces meaning that is neither entirely authored by humans nor fully automated – a co-created output characterized by distributed intentionality.

This process is neither deterministic nor random. It is probabilistically constrained and rhetorically guided. The user’s prompt initiates a direction, the model interprets it through latent patterns, and the audience completes the meaning through visual inference. The resulting message’s meaning must be reconstructed across multiple semiotic modes. Rhetorical analysis in such contexts shifts from seeking a singular authorial intent to evaluating the legitimacy of the inferences invited by the multimodal design (cf. Wu 2020). Which premises are visually implied? What values are embedded? What conclusions are made available to the audience?

Multimodality offers more than interpretive insight – it creates a critical space for ethical reflection. When persuasion operates through distributed prompts, latent aesthetics, and inferential frames, accountability can no longer be assigned to a single authorial subject. Instead, responsibility is dispersed across prompt design, the affordances and constraints of the model, and the interpretive context of the viewer. Multimodal theory thus compels us to ask: *What is being made persuasive? How is it being achieved? And to what end?*

In sum, multimodality exposes the inferential processes by which AI-generated imagery persuades through visual inference, symbolic association, and cross-modal framing. It equips us with the tools to analyze how images argue, how audiences are guided toward conclusions, and how meaning is constructed in the absence of a singular speaking subject. This inferential model of persuasion is essential for understanding AI-generated visual ideological content, where explicit claims may be absent, yet symbolic cues perform the argumentative work.

4. Ethical Theorising and the Problem of Moral Accountability

This axiological dimension underscores that generative AI is not a neutral tool but a site of ideological contestation. The values and biases embedded in prompts and training data shape not only what images are produced but also how they resonate within cultural and political contexts. Consequently, AI-generated visuals become vehicles for reinforcing or challenging dominant narratives, often in subtle and indirect ways. The rhetorical power of these images lies precisely in their ability to normalize particular worldviews by presenting them as self-evident or commonsense



through visual form. Understanding this dynamic demands a critical lens attentive to the ethical implications of prompt design, model biases, and the social frameworks that govern interpretation.

The ethical challenges posed by generative AI defy resolution through traditional models of authorship and responsibility. Conventional ethical frameworks assume a singular agent – such as a speaker, writer, or designer – whose intentions can be examined and held accountable. However, in AI-mediated communication, agency is not only obscured but also distributed, delegated, and fragmented across multiple actors: the user, the AI model, the curators of training data, and the interpretive audience. This dispersed configuration demands a new mode of ethical reasoning – one grounded in rhetorical and multimodal analysis rather than solely in computational or technical accountability (Floridi 2016).

At the core of this reconceptualization lies a simple yet profound premise: generative AI operates through language - *and language is never neutral*. Every act of classification, from interpreting a prompt to generating an image, encodes assumptions, values, and ideological positions (Crawford 2021, 123–50; Noble 2018, 171–82). Even labelling an image as “*hyperreal*”, “*feminine*”, or “*traditional*” reflects embedded cultural scripts. In rhetorical terms, this process is known as *axiologization* – the embedding of value judgments within seemingly descriptive communication (Grácio 1998, 85–98). Despite its opacity, AI-generated content actively participates in this ongoing ordering of the world.

Ethical theorizing about AI must begin by recognizing that communication itself is a form of power. It structures perception, legitimizes authority, and organizes moral hierarchies. In AI contexts, this power often operates subtly – through the aesthetics of an image, the emotional tone of a generated narrative, or the symbolic associations embedded within visual composition. These elements are not mere stylistic flourishes; they constitute ethical arguments rendered in multimodal form. To overlook this dimension is to mistake form for content – and persuasion for mere output.

A multimodal-rhetorical framework is uniquely positioned to interrogate these issues because it treats ethics not as a set of abstract principles but as a practice of meaning-making under conditions of uncertainty. It asks not only whether an image is true or fair, but how it persuades, whose worldview it reflects, and which values it promotes or erases. This approach reframes ethical evaluation from being a technical application of normative principles to a nuanced analysis of rhetorical effects. In short, it foregrounds the *how* and *why* of persuasion, not just the *what*, enabling us to ask new and critical ethical questions, such as:

1. How do AI-generated images of gendered or racialized bodies participate in the reproduction of cultural hierarchies?
2. What visual tropes are being amplified or normalized through generative media?
3. How do stylistic choices (e.g., heroic lighting, decayed architecture, ethereal glow) frame moral judgments?
4. Who benefits from these framings, and who is excluded?

These questions move the focus away from abstract principles and toward rhetorical practice. They underscore that ethical evaluation in AI cannot rely solely on fixed standards but must be situated within dialogical reasoning – where moral justification relies on the capacity to withstand critique and counterarguments.

Indeed, rhetoric is not merely a tool for persuasion; it is the very space in which values are negotiated. Ethical reasoning, from this perspective, is not grounded in universal truths but



emerges through ongoing engagement within pluralistic societies. This means debates about harm, bias, or fairness are never purely technical – they are inherently argumentative acts, shaped by power, perspective, and politics (cf. Perelman 1962; Kock 2009; Rigotti 1995; Pedersen 2015). In today's digital media landscape, these debates increasingly unfold through multimodal channels. Emotional tone, visual symbolism, facial expressions, and interface aesthetics all influence how we perceive and engage with AI. A multimodal-rhetorical approach reveals how these elements function not as mere decorative extras, but as central to the meaning and ethical force of communication.

A key contribution of this framework is its capacity to render distributed agency analyzable. Prompts are not simply technical commands – they are rhetorical events. Each prompt encodes choices about what to depict, how to depict it, and why. The user's act of prompting constitutes a form of delegated speech – a persuasive authorship enacted through the interface. Thus, the interface itself becomes a site of rhetorical affordance, providing templates that shape not only what can be said, but also what is likely to be said (Laba, 2024, 15–17).

Ultimately, we argue that AI ethics must move beyond abstract principles and embrace situated rhetorical reasoning. Ethics is not imposed from above; it is constructed through communicative interaction, cultural expectations, and symbolic form. A multimodal-rhetorical approach offers a path forward – one attentive to distributed agency, emergent meaning, and the persuasive structures through which AI systems shape social reality. In doing so, it does not claim to resolve ethical ambiguity but provides a framework for navigating it.

5. Concluding remarks

In this article, we propose that integrating rhetorical theory with multimodal analysis provides a vital framework for examining how meaning is constructed, shaped, and circulated in AI-generated content. Rather than treating AI outputs as mere technical artifacts, this approach views them as communicative events that engage in persuasive and ideological processes. Within this framework, intentionality is not conceived as a fixed mental state but as a distributed and relational function – traced through prompt creation, algorithmic mediation, and interpretive reception. Consequently, meaning-making becomes an ethically charged practice.

We also argue that developing an ethical approach to generative AI requires more than abstract normative principles; it demands attention to the discursive mechanisms through which language and imagery construct social meaning. Language models do not simply describe the world – they actively participate in defining, categorizing, and legitimizing it (Noble 2018 1–14). From this perspective, generative systems are not ethically neutral tools but rhetorical infrastructures: their outputs encode embedded assumptions, reinforce dominant narratives, and influence what is accepted as credible, persuasive, or legitimate.

At the same time, AI-generated outputs remain epistemically opaque: the pathway from prompt to image is not fully traceable. This opacity complicates causal responsibility but does not negate it. Instead, it calls for a form of forensic ethics – an effort to trace intention, effect, and persuasive structure across semiotic modes and technical systems. This entails reconstructing how meaning was produced, who influenced it, and which interpretive possibilities were activated. Žagar's theories on inferential reasoning (Žagar, 2021, 67–100) provide a valuable tool for this task – not by ascribing intent to machines, but by tracing how meaning is inferred despite the absence of traditional authorship.



We propose a way to trace intentionality not as a binary presence or absence, but as a distributed and iteratively enacted function emerging through cross-modal interaction, enabling a reconceptualization of authorship and meaning in machine-assisted communication – one that remains sensitive to ethical implications without falling into technological determinism. It calls for a shift in understanding the *locus* of meaning itself – not as the product of a singular speaker, but as a negotiated space among symbolic form, interpretive context, and agentive input, even when those agents are partly machinic. Such a perspective challenges traditional communicative paradigms by emphasizing hybridity, contingency, and the ethical labour inherent in guiding machine-mediated persuasion.

In doing so, we lay a (meta)theoretical foundation for ethical media analysis in the era of generative AI. While our argument remains conceptual, it sets the stage for future empirical research into how prompts are rhetorically constructed, how users engage in iterative meaning-making, and how multimodal outputs function persuasively within diverse socio-technical contexts. These investigations are crucial for testing and refining the framework we propose. Without such empirical grounding, ethical theorization risks becoming disconnected from the very communicative practices it seeks to illuminate.

Future research should rise to this critical challenge by bridging theoretical insights with practical investigations – examining real-world cases of AI-generated content, exploring user practices of prompt crafting and iterative refinement, and assessing the social impacts of multimodal persuasion across different communities and contexts. This work should engage interdisciplinary perspectives, combining rhetorical analysis, media studies, ethics, and human-computer interaction to develop robust, actionable frameworks. Only through such comprehensive, evidence-based inquiry can we hope to responsibly navigate the ethical complexities of shared authorship in AI-mediated communication and ensure that the evolving dynamics of meaning-making serve inclusive, just, and transparent ends.

References

- Aristotle. 2000. *Nicomachean Ethics*. Translated by Roger Crisp. Cambridge: Cambridge University Press.
- Aristotle. 2012. *The Art of Rhetoric*. London: Harper Collins.
- Anscombe, G.E.M.. 1963. *Intention*. Cambridge, MA: Harvard University Press.
- Alikhani, Malihe, Baber Khalid, and Matthew Stone. 2023. "Image-text Coherence and its Implications for Multimodal AI." *Frontiers in Artificial Intelligence* 6: e1048874.
- Barthes, Roland. 1977. *Image, Music, Text*. Translated by Stephen Heath. London: Fontana Press.
- Bender, Emily M. et al. 2021. "On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?" In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FACCT '21)*, edited by Madeleine Clare Elish, William Isaac, and Richard S. Zemel, 610–23. New York: Association for Computing Machinery.
- Bitzer, Lloyd F. 1968. "The Rhetorical Situation." *Philosophy and Rhetoric* 1 (1): 1–14.
- Buhel, Jonathan. 2016. *Assembling Arguments: Multimodal Rhetoric and Scientific Discourse*. Columbia: University of South Carolina Press.
- Bülöw, Lars, and Michael Johann. 2023. "Effects and Perception of Multimodal Recontextualization in Political Internet Memes: Evidence From Two Online Experiments in Austria." *Frontiers in Communication* 7: 1027014.
- Burke, Kenneth. 1969. *A Rhetoric of Motives*. Berkeley: University of California Press.
- Cardoso e Cunha, Tito. 2009. "Retórica da imagem?" [Rhetoric of the Image?]. In *Rhetoric and Argumentation in the Beginning of the XXist Century*, edited by Henrique Jales Ribeiro, 223–30. Coimbra: Coimbra University Press.
- Charland, Maurice. 1987. "Constitutive Rhetoric: The Case of the *Peuple Québécois*." *Quarterly Journal of Speech* 73 (2): 133–50.



- Chiang, Ted. 2024. "Why A.I. Isn't Going to Make Art". *The New Yorker*. August 31, 2024. <https://www.newyorker.com/culture/the-weekend-essay/why-ai-isnt-going-to-make-art>
- Chung, Sungmin et al. 2024. "Selective Vision Is the Challenge for Visual Reasoning: A Benchmark for Visual Argument Understanding." *arXiv*. October 23, 2024. <https://arxiv.org/abs/2406.18925>
- Crawford, Kate. *The Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence*. New Haven, CT: Yale University Press, 2021.
- Davidson, Donald. 1978. "Intending." *Philosophy of History and Action* 11: 41–60.
- Delicath, John, and Kevin Michael DeLuca. 2003. "Image Events, the Public Sphere, and Argumentative Practice: The Case of Radical Environmental Groups." *Argumentation* 17 (3): 315–33.
- Dennett, Daniel C. 1988. "The Intentional Stance in Theory and Practice." In *Machiavellian Intelligence: Social Expertise and the Evolution of Intellect in Monkeys, Apes, and Humans*, edited by Richard W. Byrne and Andrew Whiten, 180–202. Oxford: Oxford University Press.
- Fidalgo, António. "Da retórica às indústrias de persuasão." [From Rhetoric to Persuasion Industries]. In *Rhetoric and Argumentation in the Beginning of the XXIst Century*, edited by Henrique Jales Ribeiro, 232–46. Coimbra: Coimbra University Press, 2009.
- Fisher, Sarah A., Jeffrey W. Howard, and Beatriz Kira. 2024. "Moderating Synthetic Content: The Challenge of Generative AI." *Philosophy and Technology* 37: 133.
- Flewitt, Rosie, Sara Price, and Terhi Korkiakangas. 2018. "Multimodality: Methodological Explorations." *Qualitative Research* 19 (3): 146879411881741.
- Floridi, Luciano. 2016. "Faultless Responsibility: On the Nature and Allocation of Moral Responsibility for Distributed Moral Actions." *Philosophical Transactions of the Royal Society A* 374: 20160112
- Forceville, Charles. 2009. "Non-Verbal and Multimodal Metaphor in a Cognitivist Framework: Agendas for Research." In *Cognitive Linguistics: Current Applications and Future Perspectives*, edited by Gitte Kristiansene et al., 379–402. Berlin: Mouton de Gruyter.
- Gonçalves-Segundo, Pedro Rodrigues, et al. 2021. "Multimodal Argumentation: Challenges and Recent Trends. An Introduction to the Special Issue." *Revista da ABRALIN* 20 (3): 722–36.
- Grácio, Rui Alexandre. 1998. *Consequências da retórica: Para uma revalorização do múltiplo e do controverso* [Consequences of Rhetoric: Toward a Revaluation of the Multiple and the Controversial]. Coimbra: Pé de Página Editores.
- Habermas, Jürgen. *The Theory of Communicative Action, Volume 1: Reason and the Rationalization of Society*. Translated by Thomas McCarthy. Boston: Beacon Press, 1984.
- Hadley, Emily, Andrew Blatecky, and Megan Comfort. 2025. "Investigating Algorithm Review Boards for Organizational Responsible Artificial Intelligence Governance." *AI and Ethics* 5: 2485–95.
- Johnstone, Henry W.. 1990. "Rhetoric as a Wedge: A Reformulation." *Rhetoric Society Quarterly* 20 (4): 333–38.
- Kant, Immanuel. 1997. *Groundwork of the Metaphysics of Morals*. Translated by Mary Gregor. Cambridge: Cambridge University Press.
- Kjeldsen, Jens E.. 2015. "The Study of Visual and Multimodal Argumentation." *Argumentation* 29 (2): 115–32.
- Kjeldsen, Jens E.. 2016. "Symbolic Condensation and Thick Representation in Visual and Multimodal Communication." *Argumentation and Advocacy* 52 (4): 265–80.
- Kock, Christian. 2009. "Choice Is Not True or False: The Domain of Rhetorical Argumentation." *Argumentation* 23 (1): 61–80.
- Kress, Gunther R., and Theo van Leeuwen. 2001. *Multimodal Discourse: The Modes and Media of Contemporary Communication*. London: Hodder Education, 2001.
- Kuzmina, Alexandra. 2024. "Dead-End of Argumentation: The Holocaust Analogy." In *Proceedings of the Tenth Conference of the International Society for the Study of Argumentation*, edited by Ronny Boogaart et al., 519–31. Amsterdam: Sciential International Centre for Scholarship in Argumentation Theory.
- Laba, Nataliia. 2024. "Beyond Magic: Prompting for style as Affordance Actualization in Visual Generative Media." *New Media and Society*: 1–21. <https://pure.rug.nl/ws/portafiles/portal/1134456411/laba-2024-beyond-magic-prompting-for-style-as-affordance-actualization-in-visual-generative-media.pdf>
- Lendvai, Gábor F., and Gergely Gosztönyi. 2025. "Algorithmic Bias as a Core Legal Dilemma in the Age of Artificial Intelligence: Conceptual Basis and the Current State of Regulation." *Laws* 14 (13): 1–15.
- Liu, Peitong. 2020. "A Frame of Images: Intentional Reference Based on Roland Barthes' Visual Rhetoric: Take the New Media Documentary *The Human World* as an Example." *International Journal of Social Science and Education Research* 3 (12): 407–12.
- McClanahan, Bill. 2023. *Visual Criminology*. Bristol: Bristol University Press.



- Mehu, Marc. 2015. "The Integration of Emotional and Symbolic Components in Multimodal Communication." *Frontiers in Psychology* 6: 961.
- Noord, Nanne van, and Noa Garcia. 2025. "The Iconicity of the Generated Image." *arXiv*. September 19, 2025. <https://arxiv.org/abs/2509.16473>
- Noble, Safiya Umoja. 2018. *Algorithms of Oppression: How Search Engines Reinforce Racism*. New York: NYU Press.
- Pedersen, Hanne Sandberg. 2015. "Rhetorical Functions and Ethics of Narratives in Political Speeches." In *Storying Humanity: Narratives of Culture and Society*, edited by Ralf Wirth et al., 111–18. Leiden: Brill.
- Perelman, Chaïm. 1962. "Value Judgments, Justifications, and Argumentation." *Philosophy Today* 6 (1): 45–50.
- Perelman, Chaïm, and Lucie Olbrechts-Tyteca. 1969. *The New Rhetoric: A Treatise on Argumentation*. Translated by John Wilkinson and Purcell Weaver. Notre Dame: University of Notre Dame Press.
- Poulsen, Søren Vigild. 2015. "Multimodal Meaning-making." In *Key Terms in Multimodality: Definitions, Issues, Discussions*, edited by Nina Nørgaard. November 16, 2015. <https://www.sdu.dk/en/forskning/cmc/key-terms/multimodal-meaning-making>
- Pflaeging, Jana, and Hartmut Stöckl. 2021. "The Rhetoric of Multimodal Communication." *Visual Communication* 20 (3), 319–26.
- Rigotti, Francesca. 1995. "The Influence of the Rhetoric on the Ethics." *Zeitschrift für Philosophische Forschung* 49 (2): 241–58.
- Stöckl, Hartmut, and Assimakis Tsenoris. 2024. "Multimodal Rhetoric and Argumentation: Applications, Genres, Methods." *Journal of Argumentation in Context* 13 (2): 167–76.
- Trinh, Khanh, et al. 2024. "Promptly Yours? A Human Subject Study on Prompt Inference in AI-Generated Art." *arXiv*. October 10, 2024. <https://arxiv.org/abs/2410.08406>
- Wu, Tieping. 2020. "Reasoning and Appraisal in Multimodal Argumentation: Analyzing *Building a Community of a Shared Future for Mankind*." *Chinese Semiotic Studies* 16 (3): 419–38.
- Žagar, Igor Ž. 2020. "On Inference, Understanding and Interpretation in Visual Argumentation: Challenges and Problems." *Rhetoric, Visual Rhetoric and Visual Argumentation* 44 (2020): 24–53.
- Žagar, Igor Ž. 2021. *Four Critical Essays on Argumentation*. Ljubljana: Educational Research Institute.