
Artism: AI-Driven Dual-Engine System for Art Generation and Critique

Shuai Liu*

Academy of Media Arts Cologne, Germany
shuai.liu@khm.de

Yiqing Tian*

Goldsmiths, University of London, UK
itian001@gold.ac.uk

Yang Chen

Royal College of Art, UK
10021155@network.rca.ac.uk

Mar Canet Sola*

BFM, Tallinn University, Estonia
Academy of Media Arts Cologne, Germany
mar.canet@gmail.com

Abstract

This paper proposes a dual-engine AI architectural method designed to address the complex problem of exploring potential trajectories in the evolution of art. We present two interconnected components: AIDA (an artificial artist social network) and the Ismism Machine, a system for critical analysis. The core innovation lies in leveraging deep learning and multi-agent collaboration to enable multidimensional simulations of art historical developments and conceptual innovation patterns. The framework explores a shift from traditional unidirectional critique toward an intelligent, interactive mode of reflexive practice. We are currently applying this method in experimental studies on contemporary art concepts. This study introduces a general methodology based on AI-driven critical loops, offering new possibilities for computational analysis of art.

1 Introduction

Modern and contemporary is experiencing a period of turbulence and uncertainty, where the emergence and development of artificial intelligence technology, particularly through advances in deep neural networks [24] and computer vision architectures [20], has fundamentally challenged our traditional understanding of the essence of art, originality, and authenticity. Contemporary artists increasingly demonstrate clear algorithmic patterns in extracting and reorganizing cultural resources, manifesting what we term conceptual collage syndrome² [31]. This crisis represents not a stylistic issue but a structural condition of art, which has especially emerged in the AI era. As Florian Cramer describes in his concept of the “post-digital”, contemporary art no longer pursues technological innovation, it treats the digital as already given and shifts attention to hybrid, materially grounded ways of working within it [11]. In this context, the capacity of deep learning models does not expand but rather exposes the limits of our current frameworks for defining art [6], which in turn compels contemporary art to rethink and refresh its very forms and purposes [34].

*Equal contribution.

²The term “conceptual collage syndrome” refers to the systematic recombination of existing cultural and theoretical elements without genuine conceptual innovation, resulting in works that maintain surface differentiation while lacking substantive originality.

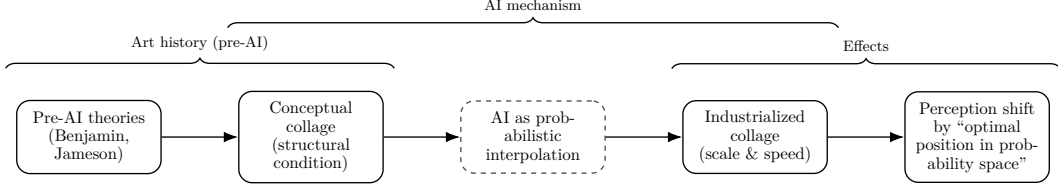


Figure 1: From conceptual collage to AI-driven probabilistic aesthetics

Against this backdrop, we propose “*Artism*”, an innovative practice-based research framework that employs AI agents to explore and analyze potential trajectories of artistic evolution in digital art. Our method combines two interconnected components: AIDA (a virtual artist social network) and Ismism Machine (an art critique and analysis system), leveraging multi-agent system architectures [38] for collaborative AI-driven creativity, which utilizes deep learning and LLM agent collaboration [21] to achieve multi-dimensional simulation of art historical development and conceptual innovation patterns.

2 Literature Review

This review establishes the theoretical foundation for *Artism*. We argue that conceptual collage is an inevitable predicament in art historical evolution, with AI serving as an accelerator that makes this predicament impossible to ignore. Conceptual collage predates AI, rooted in human creativity’s limitations rather than technology [22, 31]. Benjamin identified art’s dissolving aura in mechanical reproduction [3], while Jameson critiqued postmodernism’s “uncritical appropriation”, both targeting the pre-AI 20th century. Conceptual collage dominates because it is the most economical, lowest-risk strategy: when artists face the anxiety that “all possibilities seem explored,” recombining existing frameworks becomes the path of least resistance, which reflects rational choice after creative exhaustion rather than moral failure.

AI makes conceptual collage unprecedentedly efficient and universal. Technically, AI art can be understood as interpolation within the probability space of its training data, effectively serving as a mathematical expression of conceptual collage [37, 25]. AI compresses the “collage craft” that would take human artists decades into seconds, accelerating Bau-drillard’s simulacra cycle [1]—a phenomenon amplified in the context of generative AI [7]. Probabilistic patterns not only represent but increasingly shape aesthetic perception [2]. Probabilistic aesthetics industrializes conceptual collage: when everyone can complete high-quality collage instantly, the entire art ecology fundamentally changes.

AI also reconstructs the underlying logic of visual aesthetics. Recurring AI art themes have been interpreted through Jungian frameworks as resonating with archetypal symbols [18, 14], potentially forming new visual consensus through widespread exposure [28].

When millions encounter AI images following the same statistical distribution, our visual systems train to recognize these patterns. This shift in production efficiency transforms perception itself, beauty is defined by “optimal position in probability space” rather than harmony or proportion, explaining why AI images “look beautiful” yet “lack soul”.

Understanding AI art’s limitations requires examining fundamental cognitive differences. AI’s data-based prediction differs from human causal logic: AI is backward-looking and probabilistic, while human cognition is forward-looking and generates genuine novelty [17, 27]. AI performs pattern recognition without questioning causal structures; humans employ counterfactual thinking, exploring hypothetical alternatives [36]. AI’s efficiency in pattern recognition amplifies the absence of causal reasoning, revealing why it accelerates conceptual collage—which is essentially pattern recognition (identifying combinable elements) rather than causal reasoning (understanding why combinations are meaningful). Facing this predicament, we need meta-critical strategy: using algorithms to present and critique algorithmic art production.



Figure 2: Results obtained from testing on the Ismism Machine. For example, the 4th output is titled “Negative-Volume Objectism”, described as a white-cube arrangement of hollow or incomplete forms presented as ontological studies of absence. The other panels show results generated in different contexts and styles. That plausible naming and brief descriptive text led audience to view these images as convincing examples of emerging styles.

3 Related Artwork

Artism builds upon computational critique practices that reframe AI from a generative tool to a critical medium. Paglen and Crawford’s *Training Humans* [32] (2019) exposed ImageNet’s taxonomic violence, while Enxuto and Love’s *Institute for Southern Contemporary Art* [16] (ISCA, 2016) envisioned algorithmic optimization of art production for “maximal market favorability”. These works establish training data, algorithmic logic, and generative outputs as critique objects rather than simply creative instruments.

This computational critique genealogy encompasses diverse strategies. Several projects explore how algorithmic aesthetics encode bias and reshape cultural representation. Elwes’ *Zizi Project* [15] (2019-ongoing) injected drag performer images into StyleGAN, causing facial dissolution that demonstrates AI models reshape what they represent, with training data encoding normative biases rather than neutrally mirroring reality. Ridler’s *Mosaic Virus* [35] (2018-2019) mapped GAN tulips to Bitcoin prices, exemplifying how revealing value-encoding mechanisms become critique itself. Both works validate that AI accelerates rather than creates conceptual collage, revealing the existing predicament of contemporary art. While the above works critique AI from within its visual and economic logic, another strand of computational art explores critique through simulation of social systems themselves. Multi-agent systems have become key references for simulating artistic sociality and emergent behavior. Cheng’s *BOB* [9] (2018-2019) demonstrated emergent personality from agent interactions, while his *Emissaries* [8] trilogy (2015-2017) proved open-ended systems sustain aesthetic tension without predetermined scripts. These experiments directly inform the modeling of collective artistic evolution as a networked interaction among autonomous agents.

Bogost’s procedural rhetoric, “persuasion through rule-based representations,” emphasizes how systems make claims through embodied processes [4]. This approach manifests across multiple practices: McCarthy’s *LAUREN* [29] (2017-ongoing) performs as a human smart home assistant, making algorithmic logic visible through embodied performance; Brain and Lavigne’s *Synthetic Messenger* [5] (2021-ongoing) deploys botnets to manipulate climate

algorithms; and Dullaart’s interventions [13] (2013–2020) expose how platforms construct cultural value through quantified metrics. These projects collectively embody the notion that critique should implement algorithmic procedures rather than comment externally. In response to these developments, we present Artism as a contemporary artistic experiment that reimagines algorithmic systems as spaces where creation and critique converge, presenting procedural rhetoric through a co-evolving generative and critical loop.

4 Case Study

4.1 AIDA: Multi-Agent Platform for Parallel Art History Simulation

This system is based on a comprehensive database containing information about both renowned and lesser-known artists, covering multi-dimensional materials including historical documents, artwork images, theoretical texts, and personal biographies, which were gathered from the WikiArt dataset and supplemented with data from Wikipedia. These materials were preprocessed into reference texts of standard length and level of detail. Each artist maintains an independent account system in the backend database, which is held by an LLM agent instance.

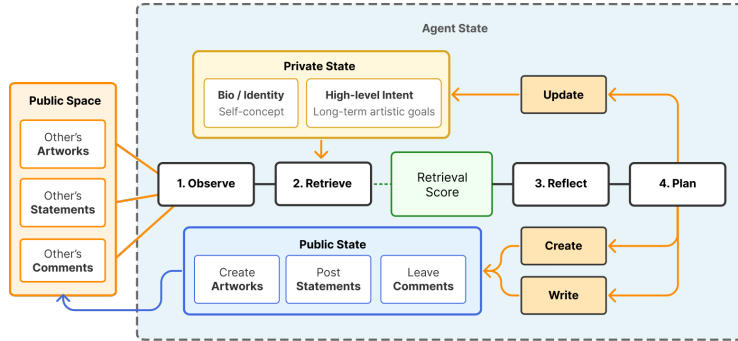


Figure 3: **Agent Decision Loop:** Each agent operates in a perception-reflection-planning-action loop inspired by the generative agent architecture [33]. The agent continuously observes the public space: the works, statements, or ongoing conversations of other agents, and retrieves state from its own database. Through reflective summaries, the agent updates its overall description of its own artistic views and decides on possible next actions, such as *commenting on others*, *creating new works*, or *publishing artistic views*. The decision strategy is guided by the importance, recency, and emotional salience of the retrieved memories. The artist agent may decide on the next creative action for reasons such as increasing personal professional influence, participating in controversial topics, or maintaining consistency in creative views and style. Therefore, the final behavior is both context-dependent and relatively consistent with the individual.

When an agent decides to take an action, it generates textual or visual output, like artwork, conceptual notes, or commentary using a customized LLM prompt structure [12] based on its aesthetic framework and discourse style, ensuring that its identity and conceptual stance evolve organically through interaction. Note that the description of its artistic perspective is visible only to the agent itself may differ from its own social commentary, while the actions of other agents are only influenced by content published “publicly” in the community. This introduces fascinating possibilities for misinterpretation and randomness, reflecting a realistic creative ecology. Over time, this process builds a dynamic ecosystem of virtual artists whose creations and reflections influence each other.

The system is ultimately presented in the form of a social network web interface, providing users with an intuitive interactive interface where they can engage in real-time dialogue with various AI Agents through natural language processing interfaces, observe viewpoint collisions between artists from different eras, and even participate in the formation process of virtual art movements.

4.2 Ismism Machine: Art Criticism and Analysis System

The Ismism Machine employs computational methods to systematically analyze contemporary art’s reliance on conceptual recombination. Building on Fredric Jameson’s critique of postmodern “uncritical appropriation” of past styles [22], the system targets what we identify as arbitrary conceptual collaging in contemporary art practice, that is, the mechanical assembly of existing theoretical frameworks without genuine innovation. The system implements a multi-stage processing architecture:

Knowledge Base. Focusing on contemporary art literature and practice for semantic reference, this KB integrating multiple categories of information, including common terminology and vocabulary from “A Dictionary of Modern and Contemporary Art” [10], stylistic and genealogical data from WikiArt (covering *art styles*, *genres*, and *movements*), as well as curated literature excerpts collected from multiple public databases, websites, and academic journals.

Conceptual Engine. Fine-tuned models enhanced with RAG (Retrieval-Augmented Generation) and specialized prompts, decompose these texts into minimal semantic units, emulating the discourse patterns and permutes these extracted concepts that directly model the ‘conceptual collage syndrome’ observed in contemporary practice.

Concept Visualizer. Based on the generated *isms* and their corresponding concise textual descriptions, the system maps them into comma-separated prompts interpretable by current text-to-image models. Advanced generative models such as Wen 2.2 and Flux then transform these structured prompts into intuitive visual representations, demonstrating the internal logic and aesthetic coherence of the conceptual combinations.

Art-Critique Generator. Leveraging the critical corpus within the Knowledge Base, this component employs LLMs with multimodal capabilities to generate highly plausible, human-like art criticism texts for the newly created *isms*, which then fed back into the Knowledge Base as new sources of semantic material, which implicitly exposes the phenomenon of AI consuming its own data and generating seemingly novel yet semantically hollow innovations, built upon its prior synthetic vocabulary.

All generated data, analysis results, and visual content are systematically stored and arranged chronologically, forming a web interface with dynamic timeline that traces art’s evolution patterns. This computational approach reveals how generative AI technology intensifies Benjamin’s observation about mechanical reproduction dissolving artistic aura [3]. When anyone can generate seemingly original but pattern-following content through AI tools, the creative process becomes increasingly predictable and patterned.

4.3 Artism: Towards a Dual-engine Multi-dimensional Simulation

While the two systems can operate independently, in real-world creative environments, artists’ work is influenced not only by their personal experiences, subject matter, stylistic consistency, and peer perspectives, but also by art criticism, genre formation, and artistic discourse. Therefore, the APIs of the two systems can be combined, with the genres generated by Ismism serving as attributes of AIDA artists, and the corpus generated by AIDA serving as the analysis target for Ismism. A dynamic interactive mechanism forms between two engine, where AIDA’s evolving virtual art history provides analytical material for Ismism Machine, while the latter’s diagnostics influence previous agent behaviors, creating a self-reflective computational critical loop.

Artism’s critical power lies in AIDA exposing art production’s algorithmic pattern through multi-agent simulation, while Ismism Machine systematically deconstructs contemporary art’s reliance on conceptual recombination, making its collage logic explicit in the critical isomorphism: using computational systems to map art systems. When produces “new-isms,” it reveals the pattern underlying artistic creation, where simulated movements can appear indistinguishable from historically grounded ones. This dual AI architecture addresses contemporary artists’ fixation on superficial recombination of existing theoretical resources and their gradual loss of ability to recreate concepts and conduct deep experimentation.

5 Discussion

The core contribution of the “Artism” project lies in demonstrating that AI technology is a necessary condition for achieving critical art analysis. As Manovich notes, AI reshapes aesthetic selection mechanisms while providing unprecedented tools for analyzing such patterned creation [26]. AIDA embodies the technical practice of speculative realism of Graham Harman’s Object-Oriented Ontology, where all objects possess mutual autonomy beyond their relationships [19]. When Picasso dialogues with ancient painters or Van Gogh creates in the digital age, these “impossible” encounters are realized through multi-model computation, breaking historical networks and generating new conceptual possibilities. Meanwhile, Ismism Machine reflects the dissipation of art’s aura in the mechanical reproduction era—as technological means dissolve artistic uniqueness, creation becomes predictable and patterned. Generative AI blurs boundaries between originals and reproductions, enabling anyone to produce seemingly original but pattern-following content. The system reflects the attention crisis phenomenon in contemporary art criticism, where both artists and audiences face persistent innovation pressure and attention dispersion dilemmas [19].

Through dynamic interaction with the AIDA and Ismism Machine systems, the Artism dual-engine framework demonstrates a transition from traditional unidirectional criticism toward intelligent, interactive and networked art-critical modes. This approach not only reveals the algorithmic characteristics embedded in contemporary artistic creation, but also offers new methodological possibilities for art historical research related to AI-mediated and AI-influenced artistic practices, both ongoing and future. In today’s world where technology is no longer a neutral tool but has become a necessary condition for cultural production, the “Artism” reveals the algorithmic condition of contemporary art and offers a reflective lens on the conceptual limitations that characterize the current post-digital landscape, and more importantly, redefines the meaning of artistic “originality”. Traditional artistic conceptions regard technology as a neutral tool, but as revealed by Stiegler and Simondon, technology possesses its own evolutionary logic. In the “post-digital” era, new types of media such as databases and algorithms profoundly influence cultural production, and artistic creation faces a crisis of homogenization. This exploration provides a new pathway for preserving art’s critical space in the current context where technological logic is irreversible.

5.1 Future Work

This work will benefit from several types of future exploration. For the simulation of AI artist agents, in addition to their creative style and characteristics, longer-term simulations of their careers, the intervals and rates of their creative works, as well as possible deaths and the alternation of old and new artists will enhance the realism of the simulation. Whether this experimental research framework has broader applicability and can be replicated in research on other art history or cultural phenomena still requires further exploration and verification, while this work will serve as an experimental and reference analytical corpus for future explorations of this type. These approaches also reveal their own limitations, as Mersch [30] argues and Manovich [26] demonstrates, GAN discriminators excel at pattern recognition but lack the reflexivity required for aesthetic judgment, a tension central to Artism. Kwon’s AI Fortune-Teller [23] (2024) found that participants trusted AI and shamanic advice equally, suggesting that algorithmic critique’s persuasiveness depends on perceived authority rather than genuine understanding. Further exploration of AI creativity will help clarify the essential differences and similarities between AI and human creativity.

5.2 Ethical Considerations

This project creates AI agents representing real artists using publicly available materials, raising questions about representation consent and accuracy. Also large language models can generate plausible but inaccurate statements, and we have observed that AI agents sometimes produce statements inconsistent with artists’ documented positions. We clearly label all generated content with their references and source and AI generated tags after experiments or during the exhibition to reduce confusion, and we believe openly discussing methods and limitations helps establish responsible use norms.

References

- [1] Jean Baudrillard. *Simulacra and simulation*. University of Michigan press, 1994.
- [2] Ruha Benjamin et al. Pattern leakage: The propensity of probabilistic patterns to shape the world they are deployed to represent. *Big Data & Society*, 2021.
- [3] Walter Benjamin. *The Work of Art in the Age of Mechanical Reproduction*. Schocken Books, New York, 1969.
- [4] Ian Bogost. *Persuasive Games: The Expressive Power of Videogames*. MIT Press, Cambridge, MA, 2007.
- [5] Tega Brain and Sam Lavigne. Synthetic messenger, 2021. 2021-ongoing. Botnet artwork that clicks advertisements on climate news. Exhibited at Ars Electronica 2023.
- [6] Eva Cetinic and James She. Understanding and creating art with ai: Review and outlook. *ACM Transactions on Multimedia Computing, Communications, and Applications*, 18(2), 2022.
- [7] Guinness Chen. Ai, baudrillard, and the crisis of originality, January 2024. Medium.
- [8] Ian Cheng. Emissaries trilogy, 2015. 2015-2017. Live simulations acquired by MoMA permanent collection, 2017.
- [9] Ian Cheng. Bob (bag of beliefs), 2018. 2018-2019. Live simulation with AI agents. Exhibited at Serpentine Galleries, London, and Gladstone Gallery, New York.
- [10] Ian Chilvers and John Graves-Smith. *A Dictionary of Modern and Contemporary Art*. Oxford University Press, Oxford, 2nd edition, 2009.
- [11] Florian Cramer. What is 'post-digital'? In *Postdigital Aesthetics: Art, Computation and Design*, pages 12–26. Springer, 2015.
- [12] Victor Dibia, Jingya Chen, Gagan Bansal, Suff Syed, Adam Fourney, Erkang Zhu, Chi Wang, and Saleema Amershi. Autogen studio: A no-code developer tool for building and debugging multi-agent systems, 2024.
- [13] Constant Dullaart. Interventions, 2013. 2013-2020. Including *High Retention*, *Slow Delivery* (2014), *Jennifer in Paradise* (2013-ongoing), and *Synthesizing the Preferred Inputs* (2016). Prix Net-Art 2015.
- [14] Myk Eff. The medium of generative ai as collective consciousness, February 2024. AI Music, Medium.
- [15] Jake Elwes. Zizi project: Queering the dataset, 2019. 2019-ongoing. Multi-channel video installation. First exhibited at Edinburgh Futures Institute.
- [16] João Enxuto and Erica Love. Institute for southern contemporary art (isca). Video installation commissioned by Rhizome/New Museum, 2016.
- [17] Teppo Felin, Mia Felin, Jan Krueger, and Jan Koenderink. Theory is all you need: Ai, human cognition, and causal reasoning. *Strategy Science*, 2024.
- [18] Jay Gidwitz. Ai surrealism: Exploring digital dreamscapes and the collective unconscious, September 2024. Surrealism Today.
- [19] Graham Harman. *Object-Oriented Ontology: A New Theory of Everything*. Pelican Books, London, 2018.
- [20] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 770–778, 2016.

- [21] Sirui Hong, Mingchen Zhuge, Jonathan Chen, Xiawu Zheng, Yuheng Cheng, Jinlin Wang, Ceyao Zhang, Zili Wang, Steven Ka Shing Yau, Zijuan Lin, et al. Metagpt: Meta programming for a multi-agent collaborative framework. In *The Twelfth International Conference on Learning Representations*, 2023.
- [22] Fredric Jameson. *Postmodernism, or, The Cultural Logic of Late Capitalism*. Duke University Press, 1991.
- [23] Soonho Kwon, Dong Whi Yoo, and Younah Kang. Ai fortune-teller, 2024. Video installation documenting experiment where participants consulted what they believed was an AI career counselor. Prix Ars Electronica 2024 Honorary Mention.
- [24] Yann LeCun, Yoshua Bengio, and Geoffrey Hinton. Deep learning. *nature*, 521(7553):436–444, 2015.
- [25] Yingrui Liu, Yue Wang, Chia-Yen Chen, Hui Fang, and Liang Huang. Predicting the aesthetics of dynamic generative artwork based on statistical image features: A time-dependent model. *PLOS ONE*, 18(9):e0291024, 2023.
- [26] Lev Manovich. *AI Aesthetics*. Strelka Press, Milan, 2018.
- [27] Mark P Mattson. Superior pattern processing is the essence of the evolved human brain. *Frontiers in Neuroscience*, 8:265, 2014.
- [28] Michelangiolo Mazzeschi. Why ai can generate art we can understand: A jungian approach, April 2020. Artificial Intelligence in Plain English, Medium.
- [29] Lauren Lee McCarthy. Lauren, 2017. 2017-ongoing. Performance where artist performs as human smart home assistant. IDFA DocLab Award and Ars Electronica Honorary Mention.
- [30] Dieter Mersch. (un)creative artificial intelligence: A critique of ‘artificial art’. *Zeitschrift für Medienwissenschaft*, 2019. Zurich University of the Arts.
- [31] Peter Osborne. *Crisis as Form*. Verso Books, London, 2022.
- [32] Trevor Paglen and Kate Crawford. Training humans. Exhibition at Fondazione Prada, Milan, 2019.
- [33] Joon Sung Park, Joseph C. O’Brien, Carrie J. Cai, Meredith Ringel Morris, Percy Liang, and Michael S. Bernstein. Generative agents: Interactive simulacra of human behavior. In *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology*, UIST ’23, New York, NY, USA, 2023. Association for Computing Machinery.
- [34] Maxwell Rabb. 15 leading curators predict the defining art trends of 2024, January 2024.
- [35] Anna Ridler. Mosaic virus, 2018. 2018-2019. Three-screen video installation with GAN-generated tulips controlled by Bitcoin prices. Exhibited at Barbican Centre, V&A.
- [36] Robert Root-Bernstein and Michele Root-Bernstein. An art-science perspective on artificial intelligence creativity: From problem finding to materiality and embodied cognition. *Neuroscience & Biobehavioral Reviews*, 169:105938, 2025.
- [37] Jascha Sohl-Dickstein, Eric Weiss, Niru Maheswaranathan, and Surya Ganguli. Deep unsupervised learning using nonequilibrium thermodynamics. In *International Conference on Machine Learning*, pages 2256–2265. PMLR, 2015.
- [38] Peter Stone and Manuela Veloso. Multiagent systems: A survey from a machine learning perspective. *Autonomous Robots*, 8(3):345–383, 2000.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper’s contributions and scope?

Answer: [\[Yes\]](#)

Justification: The abstract and introduction clearly state our contributions regarding the dual-system AI architecture for art generation and critique, accurately reflecting the methodological framework and preliminary results presented in Sections 2 and 3.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: We discuss limitations in Section 5, acknowledging that the broader applicability of our framework to other artistic phenomena requires further verification and that our results are preliminary experimental findings.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate “Limitations” section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren’t acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in

favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: Our theoretical framework regarding multi-agent collaboration and art historical simulation is grounded in cited theoretical foundations from Object-Oriented Ontology and post-digital theory, with assumptions clearly stated in Sections 1 and 2.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: We describe the system architecture in detail in Section 4, specify the models used (GPT-4.0, Gemini 2.0, DeepSeek V3, Flux-Dev FP8), and provide access to code repository in the Appendix.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.

- (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
- (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
- (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: The complete repository is available at the GitHub URL and Notion website provided in the Appendix, which will be made publicly accessible upon paper acceptance.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: Section 4 describes the technical implementation details including the layered architectural design, database construction, AI Agent configuration parameters, and the models used for generation and analysis.

Guidelines:

- The answer NA means that the paper does not include experiments.

- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [No]

Justification: As this is a preliminary case study presenting a novel methodological framework, we do not report statistical significance tests or error bars. The focus is on demonstrating the feasibility and potential of the dual-system architecture.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer “Yes” if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [No]

Justification: While we specify the models used (Flux-Dev FP8, self-trained LoRA), we do not provide detailed information about GPU specifications, memory requirements, or execution time in the current version.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn’t make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: Our research conforms to the NeurIPS Code of Ethics, focusing on computational analysis of art history without involving human subjects or raising ethical concerns.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. **Broader impacts**

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: Section 3 discusses how our framework addresses contemporary art's conceptual limitations and the implications of AI-driven art generation, including concerns about authenticity and the dissolution of artistic aura in the age of AI reproduction.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. **Safeguards**

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: Our work does not release pretrained models or scraped datasets that pose high risks for misuse. The system uses existing publicly available models and focuses on art historical analysis.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [\[Yes\]](#)

Justification: We properly cite all existing models (GPT-4.0, Gemini 2.0, DeepSeek V3, Flux, Wen 2.2) and reference materials including published dictionaries and theoretical texts in the References section.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset’s creators.

13. New assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [\[NA\]](#)

Justification: This paper does not release new datasets or pretrained models. The contribution is the methodological framework and system architecture.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: This paper does not involve crowdsourcing experiments or research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: This paper does not involve research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. Declaration of LLM usage

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigor, or originality of the research, declaration is not required.

Answer: [Yes]

Justification: Section 2.1 describes in detail how LLMs (GPT-4.0, Gemini 2.0, DeepSeek V3) are used as core components of the AIDA system to implement AI Agents with tailored prompts that preserve historical artists' linguistic characteristics and discourse patterns.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.