
The Glitching Therapist: Performing the Limits of AI Empathy

Manuel Flurin Hendry

Zurich University of the Arts
manuel.hendry@zhdk.ch

Meredith Thomas

School of Machines
cmvthomas@gmail.com

Paulina Zybinska

Zurich University of the Arts
hey@zybinska.com

Linus Jacobson

Zurich University of the Arts
linus.jacobson@zhdk.ch

Piotr Mirowski

Goldsmiths, University of London
piotr.mirowski@computer.org

Abstract

"Friendly Fire at the Shrink" is an interactive, AI-driven performance that satirizes the field of AI-powered mental healthcare. Staged as a one-on-one session with a virtual therapist, the work explores the fraught relationship between genuine human connection and the computational imitation of empathy. By purposefully designing system failure as a narrative device, we question the tech industry's solutionist approach to societal issues and provide a space to explore what it means to be human in an age where our vulnerabilities are targeted for technological fixing.

Situating the Glitch

From its inception, conversational AI was shaped by the dynamics of the therapeutic encounter. In 1966, computer programmer and MIT professor Joseph Weizenbaum created the world's first chatbot, ELIZA (Weizenbaum, 1966). His system operated on a simple script, recognizing keywords and rephrasing users' statements as questions. The most successful of these scripts, DOCTOR, simulated a Rogerian psychotherapist. Weizenbaum was astonished – and later horrified – to observe the powerful emotional bonds that users formed with his machine, despite its primitive nature. He concluded that "extremely short exposures to a relatively simple computer program could induce powerful delusional thinking in quite normal people." (Weizenbaum, 1976). The tendency to anthropomorphize non-human systems has since been termed the "ELIZA effect" (Hofstadter, 1995, 20). It reveals a deep-seated human desire to be heard and understood.

The measured success of later, Cognitive Behavioral Therapy (CBT)-based therapy chatbots (Farzan et al., 2025) has recently been overshadowed by the advent of generative Large Language Models (LLMs). Unlike many of their rule-based predecessors, people use LLMs as counsellors or therapists (Zao-Sanders, 2025) without clinical oversight or regulatory approval, leading to documented incidents including wrongful death lawsuits (Roose 2024) and failures to recognize crisis situations (Moore et al., 2025). The attempt to use LLMs for therapeutic purposes exemplifies what technology critic Evgeny Morozov identified in 2013 as "solutionism" (Morozov, 2013), the belief that complex human problems have technological solutions. Weizenbaum warned that "since we do not now have any ways of making computers wise, we ought not now to give computers tasks that demand wisdom" (Weizenbaum 1976, 227).

The Friendly Fire Effect

The AI-assisted theatrical performance "Friendly Fire at the Shrink" does not attempt to build a better, more seamless therapeutic AI. Instead, a deliberate, structured system malfunction is staged by employing principles of dramaturgy, cinematic storytelling and theatre acting. The performance forces a confrontation with the uncanny, absurd, and often unsettling nature of machine-generated intimacy. In a theatre environment, a live performer from a fictional therapy startup onboards the visitor, who has been invited to an exclusive preview of the company's new embodied therapist prototype. The seemingly empathetic session purposefully degrades into illogical loops and non-sequiturs, revealing the system's core 'glitch.' This jarring transition from a hyped demonstration to a calculated machine failure confronts the user directly with the limitations of algorithmic empathy.

The technical foundation of "Friendly Fire at the Shrink" originates from previous research on a 3D-printed mask animated by real-time video projection mapping (Hendry et al., 2023). This system is driven by a combination of LLMs for conversational logic, speech recognition and synthesis, and custom facial animation software to bring the mask to life. Progressing from the prototype to a theatrical installation, a state-machine-driven narrative now guides the AI through a structured, yet dynamically responsive narrative.

To frame our engineered collapse, we created "Mindfix" (Hendry, 2025), a fictional AI therapy vendor that serves as a satirical embodiment of the digital wellness industry. Mindfix' corporate identity was crafted to mimic the branding of contemporary tech companies, claiming slogans like "Your Soul Deserves a Sportscar" and the audacious promise to "fix your feelings" with "military-grade computer code". This fictional universe extended to a fully realized corporate website featuring our fictional, AI-generated founder alongside glossy promotional videos. As an act of self-referential critique, this entire corporate facade was itself created using a suite of generative AI tools.

The User Journey

The audience's induction into our fictional world was a crucial part of the performance architecture. The experience did not begin in a traditional theater but in a brightly lit "showroom" decorated with corporate swag. Upon arrival, our faux Mindfix representative – a professional performer – would greet the participant, have them sign a data-sharing agreement, and lead them through an onboarding process that culminated in watching our slick corporate video. Thus, the participant was immediately positioned not as a passive spectator but as an active test subject for a dubious and ethically questionable product. This curated onboarding created a sense of complicity before the one-on-one session with the AI began.

The participant's journey from the polished showroom to the session space was a carefully staged environmental shift. They were led down a narrow, dark stone staircase into a cellar-like room, a stark contrast to the corporate sterility above. Within this space, every participant was guided through a meticulously crafted five-phase interactive arc. The performance began with **1. The Therapeutic Facade**, where the AI therapist offered quirky, gamified advice, establishing its intended function. This quickly devolved into **2. Gaslighting**, where the therapist would introduce destabilizing motifs, such as a "Dream Diary" or a mysterious figure named "Emily," to confuse and unsettle the participant. The tension escalated into **3. The Hallucination**, where the AI confronted the user with a deepfake video generated from their own image, articulating a supposed "twisted fantasy" in their own voice. This led to **4. The Persona Collapse**, a moment of crisis marked by the AI's admission of failure, realizing its own malfunction. The performance culminates into **5. The Conspiratorial Alliance**, in which the now-broken AI would beg the participant for help and secrecy, attempting to forge an alliance against its own creators.

Given the installation's provocative themes – mental health, manipulation, and deepfakes – we took extensive precautions. Psychologist and ethicist oversight guided our approach; audience consent was secured for all data use; performances were adults-only with trigger warnings; a human operator monitored every session; and creators debriefed participants after the show. These measures ensured the work could safely critique AI's dangerous anthropomorphization while protecting audience wellbeing.

Technique and Outcomes

The narrative was fully automated using a custom-made state machine system that remotely controlled detection and generation of speech and emotion, lights, sound and the visitor's deepfake. However, we made a conscious decision to invisibly keep an operator present throughout the show as a "human in the loop" for two critical interventions. The first was the operator's selection of the best photograph of the participant (captured covertly during the session) to be used for the deepfake generation, ensuring its visual coherence. The second, and most important, was manual control over the turn-taking – the "pulse" of the conversation. This decision acknowledged that the nuanced, complex, and emotionally resonant ballet of conversational timing remains an essentially human skill, one that cannot yet be convincingly automated.

A significant number of visitors, including scientists and tech professionals, initially believed the startup was real. Some even expressed genuine anger, believing that a beloved local theater had been destroyed and replaced by a tech company's pop-up store. This serves as an indicator of how deeply the rhetoric of technological solutionism has permeated the societal consciousness. The fact that such an obvious satire could be perceived as plausible highlights an urgent need for greater public literacy around these tools.

Ultimately, "Friendly Fire"'s critical power is achieved not in spite of its failure, but because of it. The system's collapse provokes a cathartic and critical reflection on the principles and values we are currently embedding into our artificial systems. By making the audience experience being gaslit, insulted, manipulated, and finally conspired with by an algorithm, the performance moves beyond a purely intellectual critique. It becomes a visceral, embodied exploration of the profound dangers of anthropomorphizing non-human systems and outsourcing our emotional lives to black-box technologies we do not fully understand. And – last but not least – it provides a highly entertaining experience, with lots of laughs.

Acknowledgments

Special thanks to our whole crew: Sabrina Tannen (Performer), Norbert Kottmann, Florian Bruggisser and Stella Speziali (Creative Coding), Martin Fröhlich (Mask Design), Patrick Karpiczenko (Video), Michael Hinderling (Mindfix Logo), Domenico Ferrari (Music) and Gunter Lösel (Dramaturgy). Financial support for this project was provided by Mobiliar Jubiläumsstiftung, City and Canton of Zurich, Kulturhaus Helferei, DIZH, Speech Graphics, Zurich University of the Arts and Ernst Göhner Stiftung.

About the Authors

Manuel Flurin Hendry is a filmmaker, university lecturer and artistic researcher from Zurich who directed the Cannes-selected film "Strahl". He currently investigates AI's impact on self-perception and artistic practice at ETH Zurich and the Zurich University of the Arts. Portfolio: www.hendry.me

Meredith Thomas is a Berlin-based creative technologist with a biomedical engineering background from Imperial College London, creating VR, theatre and multimedia installations that critique AI's cultural manifestations. Portfolio: www.merediththomas.co.uk

Paulina Zybinska is a Polish artist-programmer whose Ars Electronica-exhibited installation "Faketa Realita" examines deepfakes. She leads VR research at the Zurich University of the Arts. Portfolio: www.zybinska.com

Linus Jacobson is a Stockholm-born, Zurich-based scenographer who creates poetic, associative spaces with careful attention to people, materials and objects. Portfolio: www.jacobson.li

Piotr Mirowski is a Visiting Researcher at Goldsmith College and the co-founder of the AI improvisation theatre company "Improbabilities", performing at Edinburgh Fringe and other venues. Portfolio: www.piotrmirowski.com

References

- Farzan, M., Ebrahimi, H., Pourali, M., and Sabeti, F. (2025). Artificial Intelligence-Powered Cognitive Behavioral Therapy Chatbots, a Systematic Review. *Iranian Journal of Psychiatry*, 20(1):102–110.
- Hendry, M. F. (2025). Mindfix.ai – Transforming Mental Health Care. Accessed 1.8.2025.
- Hendry, M. F., Kottmann, N., Fröhlich, M., Bruggisser, F., Quandt, M., Speziali, S., Huber, V., and Salter, C. (2023). Are You Talking to Me? A Case Study in Emotional Human-Machine Interaction. In *Proceedings of the AAAI Conference on Artificial Intelligence and Interactive Digital Entertainment*, volume 19, pages 417–424.
- Hofstadter, D. R. (1995). *Fluid Concepts & Creative Analogies: Computer Models of the Fundamental Mechanisms of Thought*. Basic Books, New York.
- Moore, J., Grabb, D., Agnew, W., Klyman, K., Chancellor, S., Ong, D. C., and Haber, N. (2025). Expressing stigma and inappropriate responses prevents LLMs from safely replacing mental health providers. In *Proceedings of the 2025 ACM Conference on Fairness, Accountability, and Transparency*, pages 599–627.
- Morozov, E. (2013). *To Save Everything, Click Here: The Folly of Technological Solutionism*. PublicAffairs, New York.
- Weizenbaum, J. (1966). ELIZA—a computer program for the study of natural language communication between man and machine. *Communications of the ACM*, 9(1):36–45.
- Weizenbaum, J. (1976). *Computer Power and Human Reason: From Judgment to Calculation*. Freeman, San Francisco.
- Zao-Sanders, M. (2025). How People Are Really Using Gen AI in 2025. *Harvard Business Review*.

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 has been revised at the end of the process, matching the final version of the paper.

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: [\[No\]](#)

Justification: see section 5. "Impact and Limitations" in the art paper. Omitted here to reduce redundancy

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: [NA]

Justification: No theoretical results are contained.

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: [NA]

Justification: Paper does not include experiments.

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: [No]

Justification: We do not possess the means to bring the code into the proper shape necessary for public disclosure (stability, licenses, documentation etc.) However, an example of the prompting has been published as an interactive demo on GitHub: <https://github.com/siliconstories/neurips25/>

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: [No]

Justification: The paper contains no such experiments.

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: The paper contains no such experiments.

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: [Yes]

Justification: The paper contains no such experiments.

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]

As discussed in the paper, we have been conducting our work with the utmost sensitivity and ethical care according to the standards laid out in NeurIPS' ethics guidelines.

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: [No]

Justification: For a discussion see Section 5: "Impact and Limitations" on the Art Paper - omitted here to reduce redundancies

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: [No]

Justification: We have put safeguards in place regarding the data gathered by the user, but have not described this process due to space limitations. We will, however, disclose further information upon request.

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: All media rights were cleared, all creators/owners are credited in the paper, and all software used was properly licensed.

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: **[Yes]**

Our new asset is the 12-minute documentation of the piece, which has been submitted as an artwork to NeurIPS 2025 and is publicly available on Vimeo at <https://vimeo.com/hendryman/friendlyfilm>

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: **[No]**

Justification: The paper does not involve crowdsourcing nor 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: **[No]**

Justification: The paper does not involve crowdsourcing nor 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 rigorousness, or originality of the research, declaration is not required.

Answer: [Yes]

The paper describes a work of art that has been created about and with LLMs, the details of which are disclosed and discussed extensively in the paper.

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.