
“Jutters”

Meike Driessen*
meike-driessen@hotmail.com

Selina Khan*
University of Amsterdam
s.j.khan@uva.nl

Gonçalo Marcelino*
University of Amsterdam
g.barretoferreiramarcelino@uva.nl

Abstract

This project explores how we engage with AI-generated content through the lens of the *jutter*: Dutch coastal foragers who comb the shoreline after storms, gathering and repurposing what the sea leaves behind. Reflecting how our lives are increasingly shaped by AI-generated media, we create a beach-like installation that blends real shoreline debris with AI-transformed images and videos. Visitors are invited to explore this space as contemporary *jutters*, deciding what to keep and what to discard. In doing so, the project reimagines AI-imagery as material for reflection, encouraging a more discerning engagement with the content that drifts through our feeds. A video preview of the installation can be found at <https://www.youtube.com/watch?v=L6319Ii7MT8>.

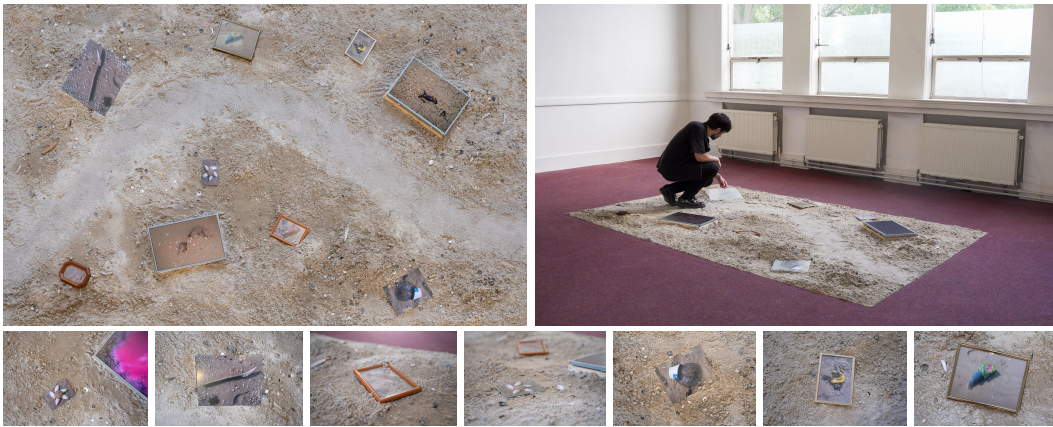


Figure 1: Overview of installation. Visitors are invited to become *jutters* and walk the sandy path to explore the artificial beach, deciding which items are worth keeping.

Background

For the longest time, in the early morning hours after a storm, people could be seen wandering the windy, grey beaches of the Netherlands. They searched for objects washed ashore, items transformed by the waves, often broken, rinsed of their history, ownership, and intended use. The Dutch language has a name for these people: *jutters*. *Jutters* were usually poor people who resided near coastal areas,

* Authors with equal contribution.

primarily living off possessions offered to them by the sea. Washed-up cargo, the wood from sunken ships, trinkets fallen overboard. These would be transformed into houses, furniture, new tools, or sold for money. Things once adrift were recontextualised, given new meaning by a tradition of curation born out of the struggle to survive [18].

Today, our messages in a bottle, our missing cargo, the things we send out, lose, or leave behind, are not only physical objects, they are also digital. They are the videos from our childhood posted to YouTube, the arguments we had in now-dead internet forums, the social media profiles we no longer have access to. Intelligent algorithms ingest and transform these digital footprints beyond recognition, and the resulting artefacts wash up on our feeds.

The majority of the content we consume digitally is no longer content we seek. It is content that is served to us, and a large part of this content is now AI-generated [23]. We are having an increasingly harder time distinguishing what is original, what is real, what is interesting, and what is purposeless [13, 14].

When we encounter an object, we engage with it not as an isolated entity but as a network of interconnected possibilities and relations [9]. It is perhaps not surprising that we are struggling to understand our relationship with AI-generated content. Like remnants of past lives washed ashore, AI-generated content lacks clear concepts of authorship, history, ownership, and intended use. The relations and possibilities that we are used to finding in things are fundamentally different or entirely absent when dealing with AI-generated content.

We aim to reframe these novel problems through an old perspective, that of the *jutter* tradition. Through this process, we hope to discover more thoughtful and effective ways of curating our digital world and unveil new networks of possibilities and relations regarding AI-generated content. With this goal in mind, we created an installation where digital images of AI-generated objects are relocated to a physical space, a small artificial beach. Participants are then invited to engage with this space, and the objects within it, with the mindset of a *jutter*. By walking through the sand and deciding, at their own pace, what is worth keeping and what can be left behind to be reclaimed by the sea, participants are encouraged to embody a crucial new type of curatorial role emerging in the age of AI-generated content.

Description of the installation

The installation is composed of an approximately 2×3 meter space covered in sand. Some parts are dry, others humid, evoking a patch of beach near the sea. Lighter-coloured sand is used to define a path through which the participant is invited to walk, with darker-coloured sand used in the remaining areas of the installation. Shells, small pieces of driftwood, and other tidal debris collected from Noordwijk beach in the Netherlands, a site once frequented by *jutters*, are scattered across the sand. Alongside these bits and pieces of a real beach are placed 7 framed prints of AI-generated images, mirroring the objects encountered by *jutters*, and 2 screens displaying AI-transformed videos, mirroring the transformation process that an object undergoes when being hit by the waves and eroded by the sea.

As participants move through the installation, they are guided by an audio narration accompanied by the sound of waves. The narrator leads the participants through the space while introducing the *jutter* tradition. The participants are invited to become *jutters* of the digital sea themselves, and to reflect on the AI-transformed objects around them. The questions posed by the narrator prompt consideration of lost histories, altered meanings, and the new roles these digitally reshaped fragments might play in future contexts. Although the installation video is delivered in Dutch with English subtitles, the narration for the installation itself is translated into English.

Technical details

To create the AI-generated images, we initially photographed over 100 objects found on the shore of Hoek van Holland beach, in the Netherlands (Figure 2). We then experimented with GPT-4 Omni [19], prompting the model to generate new images based on pairs or triplets of these photographs (Figure 3). We selected images that appeared familiar but uncanny, evoking ambiguity and inviting interpretation. By placing these convincingly real-looking, yet unfamiliar objects on the



Figure 2: A sample of the over 100 photographs of objects washed ashore taken at the beach.



Figure 3: A sample of the AI images generated based on the photographs.

artificial beach, we encouraged participants to engage in fresh acts of judgment, unanchored by prior associations.

To generate the videos, we recorded footage of objects found on the shore of Zandvoort beach in the Netherlands. We experimented with different video-to-video AI models, aiming to create videos of objects morphing while being rocked by the waves. The final videos were generated using WarpFusion [22], an AI video-to-video transformation model which distorts form and motion in a surreal, but fluid style. Stable Diffusion [21] v1.4 was used as a baseline model, since we found that earlier Diffusion models generate more surreal imagery. For post-processing, the *TopazLabs*[1] "Frame Interpolation" tool was used to increase the frame rate of the videos. Finally, minor colouring adjustments were made manually using video editing software.

Acknowledgement of video source material

The video preview incorporates imagery and audio from various external sources. Archival footage was sourced from the publicly accessible collection of the Netherlands Institute for Sound and Vision [7, 6, 5, 8], and images depicting *jutters* were obtained from the Regional Archives of Alkmaar [20, 12]. The authors took additional photographs and videos at the *Juttersmu-ZEE-um* museum, as well as in Zandvoort, Noordwijk, and Hoek van Holland. We thank the public archives for making these parts of their collections freely available, and thank the *Juttersmu-ZEE-um* for facilitating on-site photography. We also thank the *freesound.org* contributor "klankbeeld" for providing the background audio.

Biography of the authors



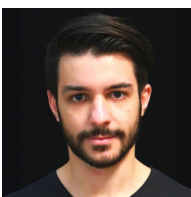
Meike Driessen

Meike is a visual artist and photographer. She graduated from St Joost School of Art & Design in 2023. In her work [4, 2, 3], she explores themes of vulnerability, loneliness, and hope. With a strong fondness for poetic visual storytelling, Meike uses photography and film as subtle, interpretive languages. Rather than offering clear answers, Meike's installations create quiet, immersive spaces where viewers can reflect, question, and connect on their terms.



Selina Khan

Selina is a research assistant at the University of Amsterdam, focused on the intersection of AI and visual arts. Her research [10, 11] focuses on capturing nuanced, context-rich information from the artistic domain in AI systems and exploring how these technologies can deepen our understanding of art while respecting its subjectivity and cultural significance. Later this year, she will start her PhD on cultural bias in multimodal AI at the University of Amsterdam.



Gonçalo Marcelino

Gonçalo is an AI researcher with a passion for art and philosophy. As part of his Master's thesis, he developed new AI methods for visual storytelling. His work led to several publications in top conferences [15, 17, 16] and received an award from the Fraunhofer Institute. He is currently a PhD candidate at the MultiX lab at the University of Amsterdam, collaborating with the Netherlands Forensic Institute to develop novel AI tools for combating money laundering.

References

- [1] Topazlabs. <https://www.topazlabs.com/>.
- [2] Meike Driessen. *bij nader inzien*. STROOM, Garage Otten Breda, 2022.
- [3] Meike Driessen. *Longing’s Embrace*. NOW SHOW, St Joost School of Art Design Breda, 2023.
- [4] Meike Driessen. *verlaat me dan als je kan*. Kunstenfestival Borkel en Schaft, 2023.
- [5] Beeld & Geluid. *Beach-life in Hoek van Holland*. 1923.
- [6] Beeld & Geluid. *Storm over de Noordzeekust*. Excerpt from *The Dutch Polygoon*, 1954.
- [7] Beeld & Geluid. *Storm raast over Nederland*. Excerpt from *The Dutch Polygoon*, 1962.
- [8] Beeld & Geluid. *Schepen op het strand*. Excerpt from *The Dutch Polygoon*, 1967.
- [9] Graham Harman. *Tool-being: Heidegger and the metaphysics of objects*. Open Court, 2011.
- [10] Selina Khan and Nanne van Noord. Stylistic multi-task analysis of ukiyo-e woodblock prints. *British Machine Vision Conference*, 2021.
- [11] Selina Khan and Nanne van Noord. Context-infused visual grounding for art. *VISArt Workshop, European Conference on Computer Vision*, 2024.
- [12] Edo Kooiman. *Texelse jeugd ruimt het strand op*. Regionaal Archief Alkmaar & Historische Vereniging Texel, 1997.
- [13] Sarah Kreps, R. McCain, and Miles Brundage. All the news that’s fit to fabricate: Ai-generated text as a tool of media misinformation. *Journal of Experimental Political Science*, 9:1–14, 11 2020.
- [14] Zeyu Lu, Di Huang, Lei Bai, Jingjing Qu, Chengyue Wu, Xihui Liu, and Wanli Ouyang. Seeing is not always believing: Benchmarking human and model perception of ai-generated images. *Advances in neural information processing systems*, 36:25435–25447, 2023.
- [15] Gonalo Marcelino, Ricardo Pinto, and Joo Magalhes. Ranking news-quality multimedia. In *Proceedings of the 2018 ACM on International Conference on Multimedia Retrieval*, pages 10–18, 2018.
- [16] Gonalo Marcelino, David Semedo, Andr Mouro, Saverio Blasi, Joo Magalhes, and Marta Mrak. Assisting news media editors with cohesive visual storylines. In *Proceedings of the 29th ACM International Conference on Multimedia*, pages 3257–3265, 2021.
- [17] Gonalo Marcelino, David Semedo, Andr Mouro, Saverio Blasi, Marta Mrak, and Joo Magalhes. A benchmark of visual storytelling in social media. In *Proceedings of the 2019 on International Conference on Multimedia Retrieval*, pages 324–328, 2019.
- [18] Jouke Minkema. *Jutters*. Pirola, Netherlands, 1988.
- [19] OpenAI. Gpt-4o system card. 2024.
- [20] L & R. *Paal 9, Frans Bakker, juten met paard en wagen voor de Meierblis*. Regionaal Archief Alkmaar & Historische Vereniging Texel, 1986.
- [21] Robin Rombach, Andreas Blattmann, Dominik Lorenz, Patrick Esser, and Bjrn Ommer. High-resolution image synthesis with latent diffusion models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 10684–10695, June 2022.
- [22] Alex Spirin. warpfusion. <https://github.com/Sxela/WarpFusion>, 2022.
- [23] Zhen Sun, Zongmin Zhang, Xinyue Shen, Ziyi Zhang, Yule Liu, Michael Backes, Yang Zhang, and Xinlei He. Are We in the AI-Generated Text World Already? Quantifying and Monitoring AIGT on Social Media. In *Annual Meeting of the Association for Computational Linguistics (ACL)*. ACL, 2025.

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 briefly describes the artwork and and thematic background of the project.

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: [\[NA\]](#)

Justification: This paper describes our installation, thematic background and its technical details and does not propose any novel technical methods.

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: We do not present theoretical results.

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: The 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: [NA]

Justification: The paper does not include experiments.

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: [NA]

Justification: The paper does not include experiments. We do provide detail on which open-source models and methods were used when applicable.

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: [NA]

Justification: The paper does not include 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: [NA]

Justification: The paper does not include 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]

Justification: We have reviewed the ethics guidelines and we confirm our work conforms to them.

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: [NA]

Justification: There is no direct societal impact associated with the work.

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: The paper poses no such risks.

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 models and codebases used in our paper are credited accordingly.

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: The paper does not release new assets.

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: Our work did 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: [NA]

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: [NA]

Justification: We do not use LLMs in such a way.

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.