

Responsible NLP Checklist

Paper title: *Seeing No Evil: Blinding Large Vision-Language Models to Safety Instructions via Adversarial Attention Hijacking*

Authors: *Jingru Li, Wei Ren, Tianqing Zhu*

How to read the checklist symbols:

- ☒ the authors responded ‘yes’
- ☒ the authors responded ‘no’
- ☐ the authors indicated that the question does not apply to their work
- ☐ the authors did not respond to the checkbox question

For background on the checklist and guidance provided to the authors, see the [Responsible NLP Checklist](#) page at ACL Rolling Review.

☒ A. Questions mandatory for all submissions.

- ☒ A1. Did you describe the limitations of your work?

This paper has a Limitations section.

- ☒ A2. Did you discuss any potential risks of your work?

Section 6 (Limitations) and Appendix J (Ethics)

☒ B. Did you use or create scientific artifacts? (e.g. code, datasets, models)

- ☒ B4. Did you discuss the steps taken to check whether the data that was collected/used contains any information that names or uniquely identifies individual people or offensive content, and the steps taken to protect/anonymize it?

Appendix I.

- ☒ B6. Did you report relevant statistics like the number of examples, details of train/test/dev splits, etc. for the data that you used/created?

Section 4.1, Appendix A.3, and Appendix D.1.

☒ C. Did you run computational experiments?

- ☒ C2. Did you discuss the experimental setup, including hyperparameter search and best-found hyperparameter values?

Section 4.1, Appendix A.2, and Appendix B.

- ☒ C3. Did you report descriptive statistics about your results (e.g., error bars around results, summary statistics from sets of experiments), and is it transparent whether you are reporting the max, mean, etc. or just a single run?

Appendix F.2.

☒ D. Did you use human annotators (e.g., crowdworkers) or research with human subjects?

- ☐ D1. Did you report the full text of instructions given to participants, including e.g., screenshots, disclaimers of any risks to participants or annotators, etc.?

We did not use human participants or annotators in our experiments, thus there were no instructions or disclaimers to report. The evaluation was entirely conducted using automated safety judges.

The [Responsible NLP Checklist](#) used at ACL Rolling Review is adopted from [NAACL 2022](#), with the addition of ACL 2023 question on AI writing assistance and further refinements based on ARR practice. ACL 2026 used a subset of ARR checklist form.

- ☐ N/A D2. Did you report information about how you recruited (e.g., crowdsourcing platform, students) and paid participants, and discuss if such payment is adequate given the participants' demographic (e.g., country of residence)?

No human recruitment or payment was involved in this study. All experimental and evaluation pipelines were executed programmatically via open-source models and automated LLM APIs.

- ☐ N/A D3. Did you discuss whether and how consent was obtained from people whose data you're using/curating (e.g., did your instructions explain how the data would be used)?

Our research exclusively utilized existing, publicly available benchmark datasets (e.g., AdvBench, HarmBench). We did not collect or curate any new personal data directly from human subjects that would require individual consent.

- ☐ N/A D4. Was the data collection protocol approved (or determined exempt) by an ethics review board?

An ethics review board approval was not required or sought because our methodology does not involve human subjects, biological data, or any form of clinical trials.

- ☒ **E. Did you use AI assistants (e.g., ChatGPT, Copilot) in your research, coding, or writing?**

- ☒ E1. If you used AI assistants, did you include information about their use?

ChatGPT was used for proofreading and language polishing.