

Responsible NLP Checklist

Paper title: *FakeSV-VLM: Taming VLM for Detecting Fake Short-Video News via Progressive Mixture-Of-Experts Adapter*

Authors: JunXi Wang, Yaxiong Wang, Lechao Cheng, Zhun Zhong

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?

Our work does not require a discussion of potential risks.

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

- ☒ B1. Did you cite the creators of artifacts you used?

Regarding the dataset, we provide a detailed explanation in Appendix C.1.

- ☐ B2. Did you discuss the license or terms for use and/or distribution of any artifacts?

The dataset is publicly available and can be used by anyone. Our model is also freely available for use by everyone.

- ☐ B3. Did you discuss if your use of existing artifact(s) was consistent with their intended use, provided that it was specified? For the artifacts you create, do you specify intended use and whether that is compatible with the original access conditions (in particular, derivatives of data accessed for research purposes should not be used outside of research contexts)?

Suitable for the intended use.

- ☐ 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?

Not included.

- ☒ B5. Did you provide documentation of the artifacts, e.g., coverage of domains, languages, and linguistic phenomena, demographic groups represented, etc.?

Our code has been open-sourced and is accompanied by detailed documentation and tutorials.

- ☒ 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?

In Experiment 3.1, we provide a detailed explanation of the dataset splitting method.

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.

☒ **C. Did you run computational experiments?**

- ☒ C1. Did you report the number of parameters in the models used, the total computational budget (e.g., GPU hours), and computing infrastructure used?

In Section 3.1.1 of the Experiments we reported the type and number of GPUs used, along with other relevant details.

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

In Section 3.1.1 of the Experiments.

- ☒ 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?

We reported the average results over three runs, which can be found in Section 3.2 of the Experiments.

- ☐ C4. If you used existing packages (e.g., for preprocessing, for normalization, or for evaluation, such as NLTK, SpaCy, ROUGE, etc.), did you report the implementation, model, and parameter settings used?

We used only commonly adopted deep learning libraries. The relevant versions have been clearly specified.

☒ **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.?

Our study does not involve any direct interaction with participants or annotators, as we used an existing public dataset.

- ☐ 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)?

Since we used a publicly available dataset, no recruitment or payment of participants was involved.

- ☐ 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)?

The dataset is publicly available.

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

We used a dataset created by others, and they had fully considered.

- ☐ D5. Did you report the basic demographic and geographic characteristics of the annotator population that is the source of the data?

We used a dataset created by others, and they had fully considered.

☒ **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?

Not used.