

A Appendix

This appendix contains the following sections:

In appendix A.1, we provide additional information on the subjects used in the fMRI data collection.

In appendix A.2, we visualize the effect of the accuracy threshold on the number of voxels we can brain-tune on.

In appendix A.3, we display the changes in brain alignment for all subjects.

In appendix A.4, we expand on the CMU MOSEI evaluation results across all emotions.

In appendix A.5, we show the STS region mask.

In appendix A.6, we discuss the potential broader impacts of our work.

In appendix A.7, we provide information on how to access our code and relevant data.

In appendix A.8, we include all licensing information.

A.1 Participants

Six healthy participants (aged 31 to 47 years at the time of recruitment in 2018), three women (sub-03, sub-04, and sub-06) and three men (sub-01, sub-02, and sub-05) were recruited to participate in the Courtois NeuroMod Project for at least 5 years. All subjects provided informed consent to participate in this study, which was approved by the ethics review board of the “CIUSS du centre-sud- de- l’île-de- Montréal” (under number CER VN 18-19-22). Three of the participants reported being native franco- phone speakers (sub- 01, sub-02, and sub-04), one as being a native anglophone (sub-06), and two as bilingual native speakers (sub-03 and sub-05). All participants reported the right hand as being their dominant hand and reported being in good general health. Exclusion criteria included visual or auditory impairments that would prevent participants from seeing and/or hearing stimuli in the scanner and major psychiatric or neurological problems. Standard exclusion criteria for MRI and MEG were also applied. Lastly, given that all stimuli and instructions are presented in English, all participants had to report having an advanced comprehension of the English language for inclusion. The above boilerplate text is taken from the cNeuroMod documentation [7], with the express intention that users should copy and paste this text into their manuscripts unchanged. It was released by the Courtois NeuroMod team under the CC0 license. For more details regarding fMRI acquisition, stimuli presentation

A.2 Cross Subject Prediction Accuracy Calculation

We follow recent studies [31], [14] in adapting [34]’s method to estimate cross-subject prediction accuracy for each voxel. For each subject, we generate all possible subsets of the remaining 5 subjects, and for each subset we use a voxel-wise encoding model (see Sec. 5) to predict one participant’s response from the others. As in previous studies [14, 31], the final value is calculated as an average at the group level. These cross-subject encoding models are trained using nine episodes (7700 TRs) from the first season of Friends, and tested on three other episodes from the same season (2872 TRs). In fig. 4, we display the number of viable voxels based on the cross subject prediction accuracy threshold. We observe that Subject-05 has no remaining voxels above 0.25, and thus deem this as our cut-off.

A.3 Differences in Normalized Brain Alignment For all Subjects

In fig. 5, we display the differences in normalized brain alignment for all subjects. Most subject models show improved alignment in and around the STS, but these improvements do not consistently extend to other regions.

A.4 CSU MOSEI Complete Emotion

In fig. 6, we breakdown the performance of the model across each emotion aggregated in the CMU MOSEI evaluation (See fig. 3, rightmost chart). Notably, sadness is the only emotion with significantly improved F1 score - however, we also observe decreased accuracy (A2). Although sadness occurs in

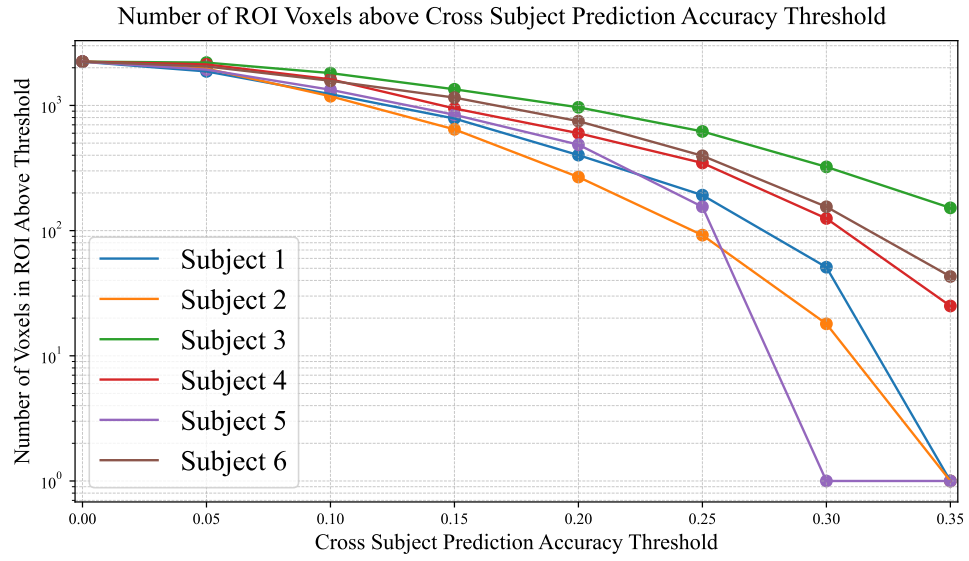


Figure 4: A subject (Subject 5) has no voxels in the STS above a cross subject prediction accuracy threshold of 0.25, and thus we cannot perform brain-tuning.

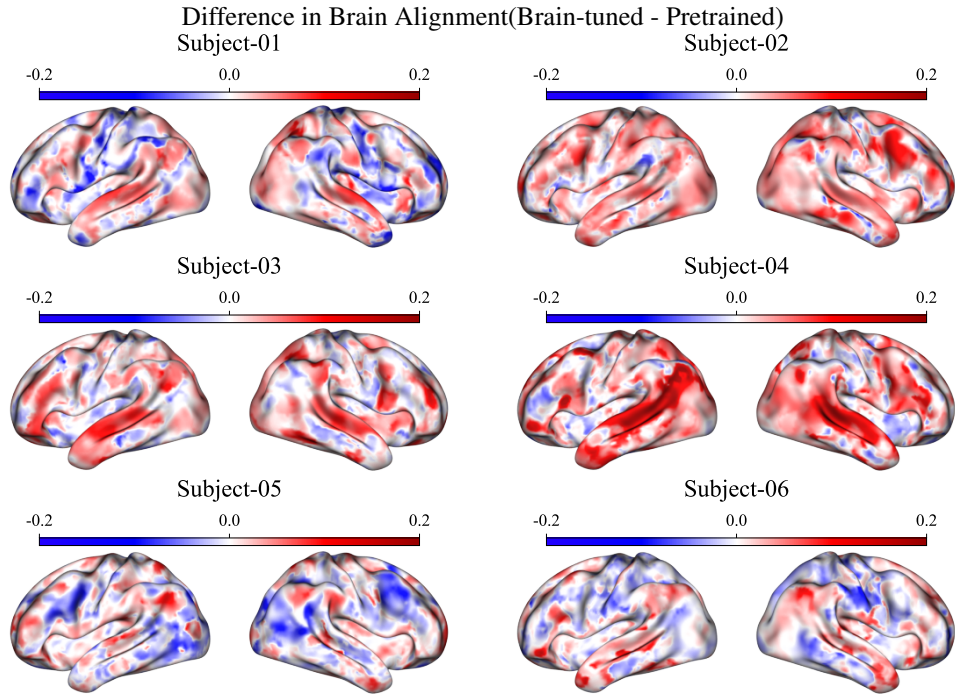


Figure 5: Differences in Normalized Brain Alignment before and after brain-tuning.

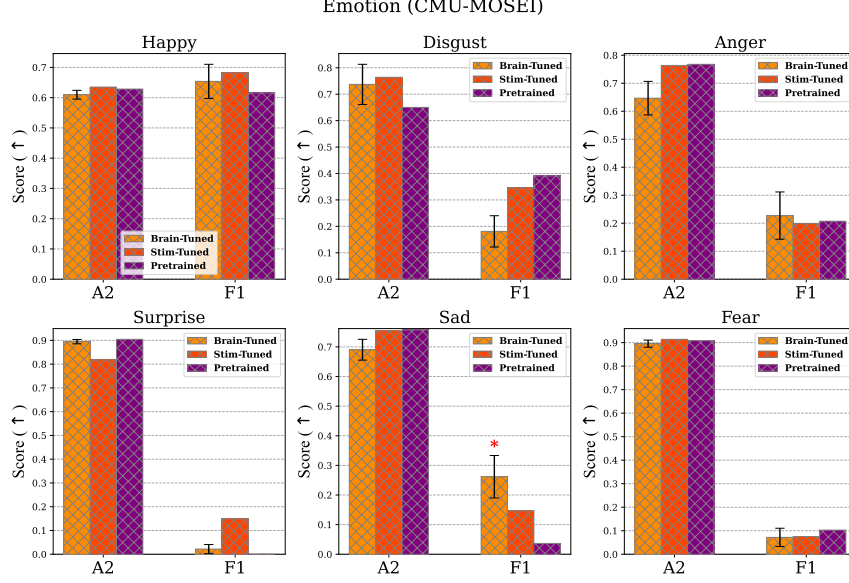


Figure 6: A breakdown of the performance of the model across each emotion aggregated in the CSU MOSEI evaluation (See fig. 3, rightmost chart).

Friends, it is not the dominant emotion in the show [33]. In future work, we hope to investigate this finding further for an explanation.

A.5 Visualization STS Region

We display the voxel mask of the STS region which we tune our model to in fig. 7.

A.6 Broader Impacts

Our work can be an initial step towards creating AI models with better understanding of human social cognition using brain activity as a tuning target. This could have positive impacts, such as improving AI-human communication or potential uses in AI-assisted therapy. Further, an AI model which can replicate human social cognition may be a useful in-silico model helping us understand social cognition in humans. On the other hand, this could enhance the abilities of AI for human manipulation. We urge future researchers to consider these pros and cons as they continue investigating this topic.

A.7 Data and Code Availability

The fMRI data used to perform the brain tuning are openly available through registered access at link <https://www.cneuromod.ca/access/access/>.

To get our code, and for exact instructions on how to replicate or results, please visit <https://huggingface.co/AnonymousSubmission43/mmbt>, download and unzip all files, and follow the instructions on the README.md in mmbt-anon.

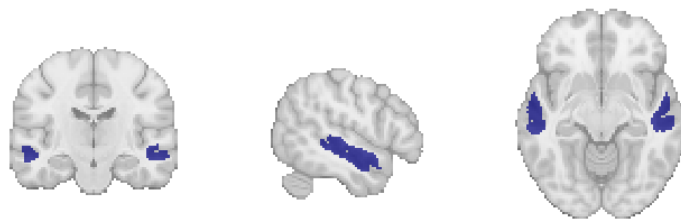
Due to privacy concerns, we do not release model weights or cross subject prediction accuracies, as these are derived from subjects' brain data.

A.8 Licenses

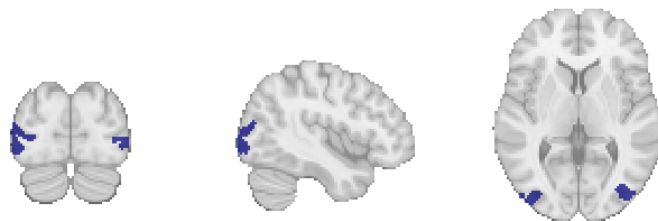
Method Diagram: Our method diagram in fig. 1 includes an audio sound wave, licensed under the public domain. You can find the url here: <https://www.pngfind.com/maxpin/bwRwR/>

Models:

STS



loc



eba

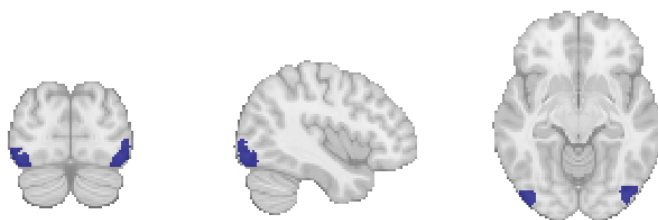


Figure 7: The STS, LOC, and EBA regions from coronal (left), lateral (middle), and horizontal (right) views.

- TVLT: MIT <https://github.com/zinengtang/TVLT/blob/main/LICENSE>

Packages:

- nilearn: BSD <https://github.com/nilearn/nilearn/blob/main/LICENSE>
- matplotlib: BSD <https://github.com/nilearn/nilearn/blob/main/LICENSE>

Datasets:

- CMU-MOSEI: MIT <https://github.com/CMU-MultiComp-Lab/CMU-MultimodalSDK?tab=MIT-1-ov-file>
- MUSTARD: MIT <https://github.com/soujanyaoria/MUSARD/blob/master/LICENSE>
- Courtois NeuroMod: CC0 <https://creativecommons.org/publicdomain/zero/1.0/legalcode>

1. Claims

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

Answer: [Yes]

Justification: In the abstract and intro, we make claims which are backed up by our experimental results - we claim increases to alignment and improvements to one of two social perception tasks, which accurately reflect the results we find through statistical testing.

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: Yes, see conclusion (section 5) for a discussion of the studies limitations (we use single model and a small number of evaluations).

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 include any theoretical results in our paper that require justification.

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: In the methods section (section 3), we fully describe our tuning architecture, the hyper parameters used during training, the specific datasets we use, how we calculate the cross subject prediction accuracies, and make all code publicly available. We do not release any weights or cross subject prediction accuracies derived from fMRI due to privacy concerns, but you can reproduce our experiments following the instructions in our paper and our code. See appendix A.7.

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: We provide open access to our code, with instructions describing how to run the main experiments in the documentation, see appendix A.7.

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: In section 3.2 we provide details on training and testing, including the data splits, optimizer type, the choice of hyperparameters, and how they were chosen. Furthermore, one can view our provided code documentation for these settings, see appendix A.7.

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: [Yes]

Justification: We carry out significance testing as described in the methods (section 3) section of the paper, for both alignment and downstream tasks. For our downstream tasks (fig. 3), we also include error bars on the brain-tuned models. These are calculated using `scipy.stats.sem` on the metric scores across our 6 brain-tuned models. We do not include error bars in our brain-alignment plots (section 3.2) - as all values are paired within participants, conventional error bars for independent samples are not applicable. We perform a paired (wilcoxon) statistical test and indicate significance in the figure instead.

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: At the end of section 3.2, we provide these details (Each brain-tuning uses 1 H100 GPU and 16 AMD EPYC 9654 CPUs on 244 GB of RAM, and takes approximately 70 hours on an H100 GPU. Each evaluation uses the same compute resources, and takes approximately 90 minutes.)

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 do not collect any data involving human participants, nor publish any novel dataset, nor expose any personal information. The CNeuromod dataset involving human subjects was collected with the approval the review board "CIUSS du centre-sud-de-l'île-de-Montréal" (under number CER VN 18-19-22).

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: In appendix A.6, we discuss potential broader impacts of our 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 (no novel data or models being released that have a possibility of misuse)

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: In appendix A.8 cite all existing assets used with their respective version.

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: [\[Yes\]](#)

Justification: We release our code, alongside details about training and evaluation of our brain-tuned model see appendix A.7.

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: [\[Yes\]](#)

Justification: We include details of stimulus presentation as is relevant to the project (subjects watched Friends half-episodes), and don't include more detail as we feel it is not necessary to understand our method or results. Subject details can be found at appendix A.1. Compensation information is not provided in the original cneuromod dataset. Further details can be found in the cneuromod documentation at <https://docs.cneuromod.ca/en/2020-alpha2/MRI.html#stimuli>

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 study did not itself collect any data from human subjects, but uses the cneuromod 2022-alpha release of the Friends dataset. All subjects provided informed consent to participate in the cneuromod data release, which was approved by the ethics review board of the “CIUSS du centre-sud-de-l’île-de-Montréal” (under number CER VN 18-19-22). For more details, see the cneuromod documentation: <https://docs.cneuromod.ca/en/latest/DATASETS.html#friends>

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: Our research does not involve any LLMs as any important, original or non-standard components.

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.