

Toward Improving Reproducibility in Neuroimaging Deep Learning Studies.

Federico Del Pup^{1,2,3*} and Manfredo Atzori^{2,3,4}

¹ Department of Information Engineering, University of Padua, Padua, Italy

² Department of Neuroscience, University of Padua, Padua, Italy

³ Padova Neuroscience Center, University of Padua, Padua, Italy

⁴ Information Systems Institute, University of Applied Sciences Western Switzerland (HES-SO Valais), Sierre, Switzerland

Correspondence*:

Federico Del Pup

federico.delpup@studenti.unipd.it

Keywords: deep learning, reproducibility, replicability, repeatability, open source, BIDS

1 INTRODUCTION

1 After more than a decade since the Imagenet breakthrough (Krizhevsky et al., 2012), there is no doubt that
2 Deep Learning (DL) has established itself as a powerful resource whose limits and risks remain difficult
3 to assess (Bengio et al., 2024). This success, originated from the ability of deep neural networks to create
4 representations of complex data with multiple levels of abstraction (LeCun et al., 2015), has inevitably
5 attracted researchers from different domains, including multidisciplinary ones.

6 In computational neuroscience (Trappenberg, 2009), deep learning has offered novel insights into the
7 functionalities of the brain, demonstrating a remarkable ability to exploit the intrinsic multimodality
8 of the field (Saxe et al., 2021). This is also reflected in the increasing volume of publications
9 encompassing diverse data types such as electroencephalographam (EEG), magnetoencephalographam
10 (MEG), structural (MRI) and functional magnetic resonance imaging (fMRI) (Zhu et al., 2019; Zhang
11 et al., 2021). Nevertheless, the potential of deep learning in neuroimaging data analysis is countered
12 by several critical issues (Miotto et al., 2018), including the scarcity of large open datasets, the poor
13 generalizability of DL models, their lack of interpretability, and the poor reproducibility of results, which
14 is discussed in this work.

15 According to (National Academies of Sciences, Engineering, and Medicine et al., 2019), reproducibility
16 is defined as the ability to “obtain consistent results using the same input data; computational steps,
17 methods, and code; and conditions of analysis.” This definition differs from that of replicability,
18 commonly mistaken as a synonym, which is instead defined as the ability to “obtain consistent results
19 across studies aimed at answering the same scientific question, each of which has obtained its own
20 data.” From the above definitions, it is possible to derive how reproducibility solely depends on how
21 authors facilitate the emulation of the same computational environment. Improving reproducibility,
22 especially in a deep learning scenario, is therefore crucial for ensuring methodological robustness and
23 result trustworthiness. However, when analyzing neuroimaging deep learning studies, several researchers
24 have expressed concerns on how authors rarely report the key elements that make a study reproducible
25 (Ciobanu-Caraus et al., 2024; Colliot et al., 2023). A recent review of DL applications for medical

image segmentation found that only 9% of the selected studies were reproducible (Renard et al., 2020), a conclusion also supported by other independent studies (Moassefi et al., 2024; Marrone et al., 2019; Ligneris et al., 2023). Furthermore, the same problem was discovered in DL-EEG applications, where, from a review of 154 selected papers, only 12 were found to be easily reproducible (Roy et al., 2019).

The above statistics highlight not only the severity of the reproducibility crisis, but also how often this important issue is overlooked by both authors and publishers. Insufficient reproducibility not only threatens the credibility of scientific findings, potentially hindering the discovery of new knowledge in the domain (e.g., treatments for neurological disorders), but also introduces inconsistencies in results due to factors such as dataset diversity, variability in preprocessing, and discrepancies in model implementation and evaluation. Such inconsistencies heighten the risk of misinterpreting data or drawing incorrect conclusions that support the validity of a framework over another, thereby posing a potential negative impact on clinical outcomes. For instance, various studies using deep learning to classify different types of dementia or to predict cognitive scores with extremely high accuracies have been questioned not only for their lack of reproducibility, but also for potential performance biases arising from inadequate data partitioning methods or ambiguous validation or model selection procedures (Brookshire et al., 2024; Wen et al., 2020). In contrast, other disciplines have achieved greater reproducibility through the use of standardized datasets and methodologies. For example, the ImageNet dataset (Krizhevsky et al., 2012) has significantly advanced computer vision by establishing benchmarks categorized by learning methods, while the General Language Understanding Evaluation (GLUE) benchmark provides a collection of resources for training, evaluating, and analyzing natural language understanding systems (Wang et al., 2018).

Consequently, there is a need for the neuroimaging field to adopt similar practices to improve the reliability of published research, which heavily depends on its reproducibility. To contribute in this direction, this paper outlines the key elements necessary for achieving reproducibility in neuroimaging deep learning studies and organize them in a table that can be also used as a checklist.

2 IMPROVING REPRODUCIBILITY

Reproducibility can only be achieved if considered since the design of the study. Even in such cases, deep learning presents numerous sources of irreproducibility, making it difficult for researchers to account for them all. Various checklists have been proposed to guide neuroscientists in designing DL studies (Roy et al., 2019; Moassefi et al., 2024), yet they often focus on specific data types or do not effectively explain how certain features affect reproducibility. To provide a clear checklist for researchers, this section introduces an updated table listing 30 key elements that are essential to ensuring reproducibility. Features in Table 1 are organized into six categories and discussed concisely in the following sub-paragraphs, highlighting how the absence of such information can hinder reproducibility and affect the reliability of results. Additionally, specific locations (paper, supplementary material, repository) are suggested to help authors improve clarity and reduce density in the papers.

2.1 Software and Hardware

Reproducibility is ensured only if authors share the source code within an open platform like GitHub¹. However, in neuroscience this is done only 10% of the times, in contrast with other domains like computer vision, where the percentage can reach 70% (Pineau et al., 2021). Sharing a well organized code is

¹ <https://github.com>

Table 1. List of criteria a research study should fulfill to ensure reproducibility.

Category		Feature	Placement
software and hardware			
	1	source code	open repository (e.g., GitHub)
	2	software environment	paper and source code
	3	computational resources	paper
dataset			
	1	raw or preprocessed data	open platform (e.g., OpenNeuro)
	2	number of subjects	paper and supplementary If multiple datasets are used, provide a summary table within the paper’s main body
	3	demographic data	
	4	number of samples	
	5	acquisition modalities	
data preprocessing			
	1	artifact handling algorithms	paper, supplementary, source code Specify the parameters used for each preprocessing step and annotate manually performed operations
	2	normalization and standardization	
	3	harmonization and co-registration	
	4	resampling	
	5	data augmentation	
model			
	1	architecture description	paper, supplementary, source code Provide a detailed summary table within the supplementary material
	2	number of learnable parameters	
	3	input dimension	
training hyperparameters			
	1	random seed	paper Provide all the information in a dedicated subsection Consider uploading the hyperparameter search results within the main repository, or within the supplementary material
	2	parameter initialization	
	3	batch size	
	4	number of epochs	
	5	loss functions	
	6	optimization algorithms	
	7	learning rate and schedulers	
	8	stopping criteria	
	9	regularization	
	10	hyperparameters search method	
model evaluation			
	1	data partition	paper and supplementary Provide the result of each training within the main repository, or as a supplementary material
	2	validation scheme	
	3	performance metric	
	4	baseline comparison	

crucial, as it provides valuable insight into the implementation correctness. Additionally, releasing an easily executable code can facilitate the design of fair comparisons between the original and new results, which helps building useful benchmarks.

On the paper side, the name and version of each main software library should be clearly reported. This helps to limit any randomness introduced by unnoticed changes in the library's code base, which can potentially compromise the reproducibility and repeatability of results (Alahmari et al., 2020). For the same reason, the entire environment should be also exported and uploaded within the repository. Furthermore, authors should report additional computational details to help tracking other known hardware non-determinism (Chen et al., 2022). These include: GPU model and number, CUDA version, training time, memory allocation, and parallelization across devices.

2.2 Dataset

Previously cited review studies agree that only a small portion of researchers share the data on dedicated open platforms. Sharing health data is indeed problematic due to strict privacy regulations. However, researchers should not feel discouraged, since there are many tools nowadays designed to facilitate data de-identification, reorganization in standardized formats (e.g., BIDS (Gorgolewski et al., 2016)), and sharing with a digital object identifier (e.g., Zenodo², OpenNeuro³ (Markiewicz et al., 2021)). Sharing raw data not only enhances the reliability of research, but also encourages other teams to try improving the original results. In addition, increasing the number of public data can facilitate the creation of multi-center datasets, which are crucial for training DL models being able to better generalize on unseen data. For example, the increasing availability large open-source datasets has encouraged researchers to explore promising deep learning architectures, such as transformers, improving the investigation of the spatiotemporal dynamics of the human brain (Kim et al., 2023; Tang et al., 2023).

Describing the dataset in detail is also important. The paper should clearly indicate the number of subjects, relevant demographic data, number of training samples, and data acquisition modalities. Data acquisition modalities may include the channel map, sampling rate, and reference for EEG data, or the scanner model, voxel size, and acquisition sequence for MRI data.

2.3 Data Preprocessing

Medical data are usually processed with complex pipelines, not necessary fully automatic. It is therefore extremely important to make this step reproducible, as different implementation of the same pipeline can lead to different results. For example, minimal changes in the ICLabel's thresholds (Pion-Tonachini et al., 2019) can drastically alter the EEG's spectral properties; similarly, the convergence threshold in the ANTs software (Avants et al., 2014) can heavily affects the MRI registration. Consequently, any customizable parameter (or manually performed operation) should be clearly described in the paper. Alternatively, researchers may upload preprocessed data alongside the raw data, enabling consistent input for analysis and assisting laboratories with limited computational resources in reproducing results. Data augmentations should also be clearly described, with particular regard to the range of random values used during their execution.

² <https://zenodo.org>

³ <https://openneuro.org>

2.4 Model architecture

Deep learning models should be described on several fronts, especially if architectures are particularly complex. First, the paper should include a schematic representation of the model, and specify its input dimension and number of trainable parameters. (If the model comes from an external library, its name and version should be reported.) Second, supplementary materials should provide a full summary table where the number of parameters, input and output dimensions, and other customizable features (e.g., weight constraints, grouped convolutions) are listed for each layer of the network. Finally, the actual model's implementation should be uploaded within the source code, as different implementations are characterized by different parameters initialization.

2.5 Training Hyperparameters

Table 1 lists a set of 10 training hyperparameters that should be described within the paper. For each of them, it is important to specify the value of each customizable parameter, if necessary in a table. Authors are encouraged to pay special attention to the use of random seeds as a way to control randomness in the code. Simply setting the random seed at the beginning of a script might not be enough. Best practice is to check if model's parameters after initialization and at the end of the training are the same on multiple rerun. Additionally, libraries like Pytorch (Paszke et al., 2019) can be configured to use a set of deterministic algorithms which, although slower, can enhance reproducibility in DL experiments.

2.6 Model evaluation

In neuroimaging data analysis, it is best practice to evaluate DL models with subject-based cross-validation procedures, which require to partition the dataset multiple times in training and validation sets (Kunjan et al., 2021). Partitions should be reproducible, because different splits can dramatically change the results due to strong subject-based characteristics that can guide the learning process towards different loss minima. Furthermore, it is important to clearly state how the validation set is used during training. While switching to a train-validation-test split scenario as part of a nested cross-validation, as proposed in (Kim et al., 2022), is preferable, it is unfortunately usual to find works that use the validation set to stop the training. This process introduces data leakage and affects the reliability of the results. Nevertheless, this information is often omitted (Pandey and Seeja, 2022), leaving the question of whether data leakage was introduced unanswered.

3 DISCUSSION

Designing a reproducible deep learning experiment, especially in a complex and heterogeneous field like neuroscience, is extremely challenging. Researchers must consider several factors that can influence experimental outcomes; otherwise, reproducibility can be permanently lost. In computational neuroscience, DL applications to neuroimaging data require multiple training instances to be run as part of subject-based cross-validation procedures (e.g., Leave-N-Subject-Out, typically using 5 or 10 folds (Poldrack et al., 2020)). Even with modern GPUs capable of performing 10^{12} floating point operations per second, these strategies can be time-consuming and computationally intensive. It is natural that researchers might hesitate to re-run the analysis if they encounter reproducibility issues. Additionally, deep learning advances rapidly, with numerous papers published every month. While this rapidity fuels the research teams' need to quickly disseminate their work, it should also not undermine reproducibility, which is central to the scientific method and a necessary (though not sufficient) condition for a scientific

statement to be accepted as new knowledge (Colliot et al., 2023). Focusing on reproducibility indirectly compels researchers to develop methodologically robust experiments that yield more reliable results.

Technical reproducibility in DL studies can only be assured if both the code and dataset are released by the study group. Although privacy regulations can complicate data sharing, teams can nonetheless make their source code publicly available (Ciobanu-Caraus et al., 2024). This will provide a complementary resource that is often necessary for clearly describing neuroimaging deep learning studies, especially when articles are subjected to word or page limits. Table 1 not only lists 30 key reproducibility elements, but also suggests specific places to present them without hindering the paper's readability. However, this alone might not be enough. Actions are required also from the publisher side; in fact, many of them have begun revising their policies to strongly encourage, and occasionally mandate, the sharing of code and research data.

In conclusion, while deep learning is a powerful tool, it still remains extremely sensitive to minimal variations in the experimental setup. It is essential to re-evaluate priorities, especially in sensitive applications like medical ones, and look beyond mere model performance. Without addressing critical issues such as reproducibility and generalizability, reported accuracies may be seen as mere numbers lacking practical meaning.

AUTHOR CONTRIBUTIONS

FDP: Conceptualization, Writing – original draft, Writing – review & editing. MA: Supervision, Funding acquisition, Writing – review & editing.

FUNDING

The author(s) declare financial support was received for the research, authorship, and/or publication of this article. This work was supported by the European Union's Horizon Europe research and innovation programme under Grant agreement no 101137074 - HEREDITARY. This work has received funding STARS@UNIPD funding program of the University of Padova, Italy, through the project: MEDMAX.

ACKNOWLEDGMENTS

FDP would like to thank the other members of the Padova Neuroscience Center for their inspiring comments on the topic presented in this paper.

CONFLICT OF INTEREST

The authors declare that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.

REFERENCES

- Alahmari, S. S., Goldgof, D. B., Mouton, P. R., and Hall, L. O. (2020). Challenges for the repeatability of deep learning models. *IEEE Access* 8, 211860–211868. doi:10.1109/ACCESS.2020.3039833
- Avants, B. B., Tustison, N. J., Stauffer, M., Song, G., Wu, B., and Gee, J. C. (2014). The insight toolkit image registration framework. *Frontiers in Neuroinformatics* 8. doi:10.3389/fninf.2014.00044

- 171 Bengio, Y., Hinton, G., Yao, A., Song, D., Abbeel, P., Darrell, T., et al. (2024). Managing extreme AI
172 risks amid rapid progress. *Science* 384, 842–845. doi:10.1126/science.adn0117
- 173 Brookshire, G., Kasper, J., Blauch, N. M., Wu, Y. C., Glatt, R., Merrill, D. A., et al. (2024). Data leakage
174 in deep learning studies of translational EEG. *Frontiers in Neuroscience* 18. doi:10.3389/fnins.2024.
175 1373515
- 176 Chen, B., Wen, M., Shi, Y., Lin, D., Rajbahadur, G. K., and Jiang, Z. M. J. (2022). Towards training
177 reproducible deep learning models. In *Proceedings of the 44th International Conference on Software
178 Engineering* (New York, NY, USA: Association for Computing Machinery), ICSE '22, 2202–2214.
179 doi:10.1145/3510003.3510163
- 180 Ciobanu-Caraus, O., Aicher, A., Kernbach, J. M., Regli, L., Serra, C., and Staartjes, V. E. (2024). A
181 critical moment in machine learning in medicine: on reproducible and interpretable learning. *Acta
182 Neurochirurgica* 166, 14. doi:10.1007/s00701-024-05892-8
- 183 Colliot, O., Thibeu-Sutre, E., and Burgos, N. (2023). *Reproducibility in Machine Learning for Medical
184 Imaging* (New York, NY: Springer US), chap. 3. 631–653. doi:10.1007/978-1-0716-3195-9_21
- 185 Gorgolewski, K. J., Auer, T., Calhoun, V. D., Craddock, R. C., Das, S., Duff, E. P., et al. (2016). The brain
186 imaging data structure, a format for organizing and describing outputs of neuroimaging experiments.
187 *Scientific data* 3, 1–9. doi:10.1038/sdata.2016.44
- 188 Kim, J., Hwang, D.-U., Son, E. J., Oh, S. H., Kim, W., Kim, Y., et al. (2022). Emotion recognition while
189 applying cosmetic cream using deep learning from EEG data; cross-subject analysis. *PLOS ONE* 17,
190 1–26. doi:10.1371/journal.pone.0274203
- 191 Kim, P., Kwon, J., Joo, S., Bae, S., Lee, D., Jung, Y., et al. (2023). Swift: Swin 4d fmri transformer.
192 *Advances in Neural Information Processing Systems* 36, 42015–42037
- 193 Krizhevsky, A., Sutskever, I., and Hinton, G. E. (2012). Imagenet classification with deep convolutional
194 neural networks. In *Advances in Neural Information Processing Systems*, eds. F. Pereira, C. Burges,
195 L. Bottou, and K. Weinberger (Curran Associates, Inc.), vol. 25, . doi:10.1145/3065386
- 196 Kunjan, S., Grummet, T. S., Pope, K. J., Powers, D. M., Fitzgibbon, S. P., Bastiampillai, T., et al.
197 (2021). The necessity of leave one subject out (LOSO) cross validation for EEG disease diagnosis.
198 In *Brain Informatics: 14th International Conference, BI 2021, Virtual Event, September 17–19, 2021,
199 Proceedings 14* (Springer), 558–567. doi:10.1007/978-3-030-86993-9_50
- 200 LeCun, Y., Bengio, Y., and Hinton, G. (2015). Deep learning. *nature* 521, 436–444. doi:10.1038/
201 nature14539
- 202 Ligneris, M. D., Bonnet, A., Chatelain, Y., Glatard, T., Sdika, M., Vila, G., et al. (2023). Reproducibility
203 of tumor segmentation outcomes with a deep learning model. In *2023 IEEE 20th International
204 Symposium on Biomedical Imaging (ISBI)*. 1–5. doi:10.1109/ISBI53787.2023.10230482
- 205 Markiewicz, C. J., Gorgolewski, K. J., Feingold, F., Blair, R., Halchenko, Y. O., Miller, E., et al. (2021).
206 The OpenNeuro resource for sharing of neuroscience data. *Elife* 10, e71774
- 207 Marrone, S., Olivieri, S., Piantadosi, G., and Sansone, C. (2019). Reproducibility of deep CNN for
208 biomedical image processing across frameworks and architectures. In *2019 27th european signal
209 processing conference (eusipco)* (IEEE), 1–5. doi:10.23919/EUSIPCO.2019.8902690
- 210 Miotto, R., Wang, F., Wang, S., Jiang, X., and Dudley, J. T. (2018). Deep learning for healthcare: review,
211 opportunities and challenges. *Briefings in bioinformatics* 19, 1236–1246. doi:10.1093/bib/bbx044
- 212 Moassefi, M., Singh, Y., Conte, G. M., Khosravi, B., Rouzrokh, P., Vahdati, S., et al. (2024). Checklist
213 for reproducibility of deep learning in medical imaging. *Journal of Imaging Informatics in Medicine* ,
214 1–10doi:10.1007/s10278-024-01065-2

- 215 National Academies of Sciences, Engineering, and Medicine, Policy and Global Affairs, Committee on
216 Science, Engineering, Medicine, and Public Policy, Board on Research Data and Information, Division
217 on Engineering and Physical Sciences, Committee on Applied and Theoretical Statistics, et al. (2019).
218 *Reproducibility and replicability in science* (National Academies Press). doi:10.17226/25303
- 219 Pandey, P. and Seeja, K. (2022). Subject independent emotion recognition from EEG using VMD and
220 deep learning. *Journal of King Saud University - Computer and Information Sciences* 34, 1730–1738.
221 doi:10.1016/j.jksuci.2019.11.003
- 222 Paszke, A., Gross, S., Massa, F., Lerer, A., Bradbury, J., Chanan, G., et al. (2019). PyTorch: an imperative
223 style, high-performance deep learning library. In *Proceedings of the 33rd International Conference on*
224 *Neural Information Processing Systems* (Red Hook, NY, USA: Curran Associates Inc.)
- 225 Pineau, J., Vincent-Lamarre, P., Sinha, K., Lariviere, V., Beygelzimer, A., d'Alche Buc, F., et al. (2021).
226 Improving reproducibility in machine learning research (a report from the NeurIPS 2019 reproducibility
227 program). *Journal of Machine Learning Research* 22, 1–20
- 228 Pion-Tonachini, L., Kreutz-Delgado, K., and Makeig, S. (2019). ICLabel: An automated
229 electroencephalographic independent component classifier, dataset, and website. *NeuroImage* 198,
230 181–197. doi:10.1016/j.neuroimage.2019.05.026
- 231 Poldrack, R. A., Huckins, G., and Varoquaux, G. (2020). Establishment of best practices for evidence for
232 prediction: A review. *JAMA Psychiatry* 77, 534–540. doi:10.1001/jamapsychiatry.2019.3671
- 233 Renard, F., Guedria, S., Palma, N. D., and Vuillerme, N. (2020). Variability and reproducibility
234 in deep learning for medical image segmentation. *Scientific Reports* 10, 13724. doi:10.1038/
235 s41598-020-69920-0
- 236 Roy, Y., Banville, H., Albuquerque, I., Gramfort, A., Falk, T. H., and Faubert, J. (2019). Deep
237 learning-based electroencephalography analysis: a systematic review. *Journal of Neural Engineering*
238 16, 051001. doi:10.1088/1741-2552/ab260c
- 239 Saxe, A., Nelli, S., and Summerfield, C. (2021). If deep learning is the answer, what is the question?
240 *Nature Reviews Neuroscience* 22, 55–67. doi:10.1038/s41583-020-00395-8
- 241 Tang, C., Wei, M., Sun, J., Wang, S., and Zhang, Y. (2023). CsAGP: Detecting alzheimer's disease from
242 multimodal images via dual-transformer with cross-attention and graph pooling. *Journal of King Saud*
243 *University - Computer and Information Sciences* 35, 101618. doi:10.1016/j.jksuci.2023.101618
- 244 Trappenberg, T. P. (2009). *Fundamentals of Computational Neuroscience* (Oxford University Press).
245 doi:10.1093/oso/9780199568413.001.0001
- 246 Wang, A., Singh, A., Michael, J., Hill, F., Levy, O., and Bowman, S. (2018). GLUE: A multi-task
247 benchmark and analysis platform for natural language understanding. In *Proceedings of the 2018*
248 *EMNLP Workshop BlackboxNLP: Analyzing and Interpreting Neural Networks for NLP*, eds. T. Linzen,
249 G. Chrupała, and A. Alishahi (Brussels, Belgium: Association for Computational Linguistics), 353–
250 355. doi:10.18653/v1/W18-5446
- 251 Wen, J., Thibeau-Sutre, E., Diaz-Melo, M., Samper-González, J., Routier, A., Bottani, S., et al. (2020).
252 Convolutional neural networks for classification of Alzheimer's disease: Overview and reproducible
253 evaluation. *Medical Image Analysis* 63, 101694. doi:10.1016/j.media.2020.101694
- 254 Zhang, X., Yao, L., Wang, X., Monaghan, J., Mcalpine, D., and Zhang, Y. (2021). A survey on
255 deep learning-based non-invasive brain signals: recent advances and new frontiers. *Journal of neural*
256 *engineering* 18, 031002. doi:10.1088/1741-2552/abc902
- 257 Zhu, G., Jiang, B., Tong, L., Xie, Y., Zaharchuk, G., and Wintermark, M. (2019). Applications of deep
258 learning to neuro-imaging techniques. *Frontiers in neurology* 10, 869. doi:10.3389/fneur.2019.00869