

RESEARCH ARTICLE

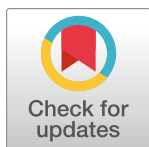
Numerical uncertainty in analytical pipelines lead to impactful variability in brain networks

Gregory Kiar^{1*}, Yohan Chatelain², Pablo de Oliveira Castro³, Eric Petit⁴, Ariel Rokem⁵, Gaël Varoquaux⁶, Bratislav Misic¹, Alan C. Evans¹, Tristan Glatard²

1 Montréal Neurological Institute, McGill University, Montréal, QC, Canada, **2** Department of Computer Science and Software Engineering, Concordia University, Montréal, QC, Canada, **3** Department of Computer Science, Université de Versailles, Versailles, France, **4** Exascale Computing Lab, Intel, Paris, France, **5** Department of Psychology and eScience Institute, University of Washington, Seattle, WA, United States of America, **6** Parietal Project-team, INRIA Saclay-ile de France, Paris, France

✉ These authors contributed equally to this work.

* gregory.kiar@childmind.org



OPEN ACCESS

Citation: Kiar G, Chatelain Y, de Oliveira Castro P, Petit E, Rokem A, Varoquaux G, et al. (2021) Numerical uncertainty in analytical pipelines lead to impactful variability in brain networks. PLoS ONE 16(11): e0250755. <https://doi.org/10.1371/journal.pone.0250755>

Editor: Stavros I. Dimitriadis, Cardiff University, UNITED KINGDOM

Received: April 19, 2021

Accepted: August 25, 2021

Published: November 1, 2021

Peer Review History: PLOS recognizes the benefits of transparency in the peer review process; therefore, we enable the publication of all of the content of peer review and author responses alongside final, published articles. The editorial history of this article is available here: <https://doi.org/10.1371/journal.pone.0250755>

Copyright: © 2021 Kiar et al. This is an open access article distributed under the terms of the [Creative Commons Attribution License](https://creativecommons.org/licenses/by/4.0/), which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

Data Availability Statement: The unprocessed dataset is available through The Consortium of Reliability and Reproducibility (http://fcon_1000.projects.nitrc.org/indi/enhanced/), including both

Abstract

The analysis of brain-imaging data requires complex processing pipelines to support findings on brain function or pathologies. Recent work has shown that variability in analytical decisions, small amounts of noise, or computational environments can lead to substantial differences in the results, endangering the trust in conclusions. We explored the instability of results by instrumenting a structural connectome estimation pipeline with Monte Carlo Arithmetic to introduce random noise throughout. We evaluated the reliability of the connectomes, the robustness of their features, and the eventual impact on analysis. The stability of results was found to range from perfectly stable (i.e. all digits of data significant) to highly unstable (i.e. 0 – 1 significant digits). This paper highlights the potential of leveraging induced variance in estimates of brain connectivity to reduce the bias in networks without compromising reliability, alongside increasing the robustness and potential upper-bound of their applications in the classification of individual differences. We demonstrate that stability evaluations are necessary for understanding error inherent to brain imaging experiments, and how numerical analysis can be applied to typical analytical workflows both in brain imaging and other domains of computational sciences, as the techniques used were data and context agnostic and globally relevant. Overall, while the extreme variability in results due to analytical instabilities could severely hamper our understanding of brain organization, it also affords us the opportunity to increase the robustness of findings.

Introduction

The modelling of brain networks, called connectomics, has shaped our understanding of the structure and function of the brain across a variety of organisms and scales over the last decade [1–6]. In humans, these wiring diagrams are obtained *in vivo* through Magnetic Resonance Imaging (MRI), and show promise towards identifying biomarkers of disease. This can not

the imaging data as well as phenotypic data which may be obtained upon submission and compliance with a Data Usage Agreement. The connectomes generated through simulations have been bundled and stored permanently (<https://doi.org/10.5281/zenodo.4041549>), and are made available through The Canadian Open Neuroscience Platform (<https://portal.conp.ca/search>, search term “Kiar”). All software developed for processing or evaluation is publicly available on GitHub at <https://github.com/gkpapers/2020ImpactOfInstability>. Experiments were launched using Boutiques and Clowdr in Compute Canada’s HPC cluster environment. MCA instrumentation was achieved through Verificarlo available on Github at <https://github.com/verificarlo/verificarlo>. A set of MCA instrumented software containers is available on Github at <https://github.com/gkiar/fuzzy>.

Funding: GK was fully supported for this work by the Natural Sciences and Engineering Research Council of Canada <https://www.nserc-crsng.gc.ca/index_eng.asp> Award number: CGSD3 - 519497 - 2018. The funders had no role in study design, data collection and analysis, decision to publish, or preparation of the manuscript.

Competing interests: The authors have declared that no competing interests exist.

only improve understanding of so-called “connectopathies”, such as Alzheimer’s Disease and Schizophrenia, but potentially pave the way for therapeutics [7–11].

However, the analysis of brain imaging data relies on complex computational methods and software. Tools are trusted to perform everything from pre-processing tasks to downstream statistical evaluation. While these tools undoubtedly undergo rigorous evaluation on bespoke datasets, in the absence of ground-truth this is often evaluated through measures of reliability [12–16], proxy outcome statistics, or agreement with existing theory. Importantly, this means that tools and datasets are not necessarily of known or consistent quality, and it is not uncommon that equivalent experiments may lead to diverging conclusions [17–23]. While many scientific disciplines suffer from a lack of reproducibility [24], this was recently explored in brain imaging by a 70 team consortium which performed equivalent analyses and found widely inconsistent results [17], and it is likely that software instabilities played a role. This study does not broach dataset differences, but there are considerable works which demonstrate that data selection may compound these effects (e.g. [14, 16]).

The present study approached evaluating reproducibility from a computational perspective in which a series of brain imaging studies were numerically perturbed in such a way that the plausibility of results was not affected, and the implications of the observed instabilities on downstream analyses were quantified. We accomplished this through the use of Monte Carlo Arithmetic (MCA) [25, 26], a technique which enables characterization of the sensitivity of a system to small numerical perturbations. This is importantly distinct from data perturbation experiments where the underlying datasets are manipulated or pathologies may be simulated, and allows for the evaluation of experimental uncertainty in real-world settings. We explored the impact of numerical perturbations through the direct comparison of structural connectomes, the consistency of their features, and their eventual application in a neuroscience study. We also characterized the consequences of instability in these pipelines on the reliability of derived datasets, and discuss how the induced variability may be harnessed to increase the discriminability of datasets, in an approach akin to ensemble learning. Finally, we make recommendations for the roles perturbation analyses may play in brain imaging research and beyond.

Results

Graphs vary widely with perturbations

Prior to exploring the analytic impact of instabilities, a direct understanding of the induced variability was required. A subset of the Nathan Kline Institute Rockland Sample (NKIRS) dataset [27] was randomly selected to contain 25 individuals with two sessions of imaging data, each of which was subsampled into two components, resulting in four samples per individual and 100 samples total ($25 \times 2 \times 2$ samples). Structural connectomes were generated with canonical deterministic and probabilistic pipelines [28, 29] which were instrumented with MCA, replicating computational noise either sparsely or densely throughout the pipelines [19, 26]. In the sparse case, a small subset of the libraries were instrumented with MCA, allowing for the evaluation of the cascading effects of numerical instabilities that may arise. In the dense case, operations are more uniformly perturbed and thus the law of large numbers suggests that perturbations will quickly offset one-another and only dramatic local instabilities will have propagating effects. Importantly, the perturbations resulting from the sparse setting represent a strict subset of the possible outcomes of the dense implementation. The random perturbations are statistically independent from one another across both settings and simulations. Instrumenting pipelines with MCA increases their computation time, in this case by multiplication factors of $1.2 \times$ and $7 \times$ for the sparse and dense settings, respectively [19]. The

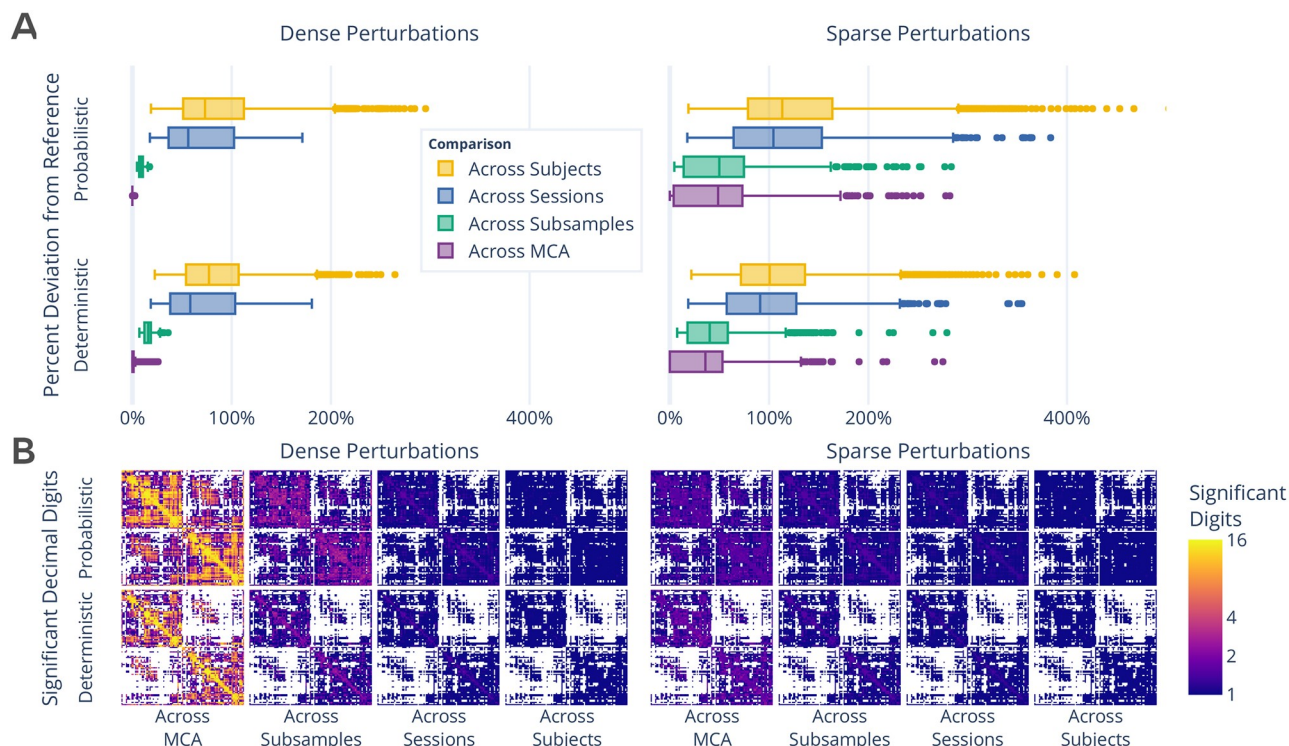


Fig 1. Exploration of perturbation-induced deviations from reference structural connectomes. (A) The absolute deviations between connectomes, in the form of normalized percent deviation from reference. The difference in MCA-perturbed connectomes is shown as the across MCA series, and is presented relative to the variability observed across subsamples, sessions, and subjects. (B) The number of significant decimal digits in each set of connectomes as obtained by evaluating the complete distribution of networks. In the case of 16, values can be fully relied upon, whereas in the case of 1 only the first digit of a value can be trusted. Dense and sparse perturbations are shown on the left and right, respectively.

<https://doi.org/10.1371/journal.pone.0250755.g001>

results obtained were compared to unperturbed (e.g. reference) connectomes in both cases. The connectomes were sampled 20 times per sample and once without perturbations, resulting in a total of 8,400 connectomes. Two versions of the unperturbed connectomes were generated and compared such that the absence of variability aside from that induced via MCA could be confirmed.

The stability of structural connectomes was evaluated through the normalized percent deviation from reference [19] and the number of significant digits (Fig 1). The comparisons were grouped according to differences across simulations, subsampling of data, sessions of acquisition, or subjects, and accordingly sorted from most to least similar. While the similarity of connectomes decreases as the collections become more distinct, connectomes generated with sparse perturbations show considerable variability, often reaching deviations equal to or greater than those observed across individuals or sessions (Fig 1A; right). Interpreting these results with respect to the distinct MCA environments used suggests that the tested pipelines may not suffer from single dominant sources of instability, but that nevertheless there exist minor local instabilities which may propagate throughout the pipeline. Furthermore, this finding suggests that instabilities inherent to these pipelines may mask session or individual differences, limiting the trustworthiness of derived connectomes. While both pipelines show similar performance, the probabilistic pipeline was more stable in the face of dense perturbations whereas the deterministic was more stable to sparse perturbations ($p < 0.0001$ for all;

exploratory). As an alternative to the normalized percent deviation, the stability of correlations between networks can be found in S1 Section in [S1 File](#).

The number of significant digits per edge across connectomes ([Fig 1B](#)) similarly decreases alongside the decreasing similarity between comparison groups. While the cross-MCA comparison of connectomes generated with dense perturbations show nearly perfect precision for many edges (approaching the maximum of 15.7 digits for 64-bit data), this evaluation uniquely shows considerable drop off in performance when comparing networks across subsamplings (average of <4 digits). In addition, sparsely perturbed connectomes show no more than an average of 3 significant digits across all comparison groups, demonstrating a significant limitation in the reliability of independent edge weights. The number of significant digits across individuals did not exceed a single digit per edge in any case, indicating that only the order of magnitude of edges in naively computed groupwise average connectomes can be trusted. The combination of these results with those presented in [Fig 1A](#) suggests that while specific edge weights are largely affected by instabilities, macro-scale network structure is stable.

Sparse perturbations reduce off-target signal

We assessed the reproducibility of the dataset through mimicking and extending a typical test-retest experiment [[14](#)] in which the similarity of samples across sessions were compared to distinct samples in the dataset ([Table 1](#), with additional experiments and explanation of the measure and its scaling in S2 Section in [S1 File](#)). The ability to discriminate connectomes across subjects (Hypothesis 1) is an essential prerequisite for the application of brain imaging towards identifying individual differences [[6](#)]. In testing hypothesis 1, we observe that the dataset is discriminable with a scaled score of 0.82 ($p < 0.001$; optimal score: 1.0; chance: 0.04) for both pipelines in the absence of MCA. We can see that inducing instabilities through MCA preserves the discriminability in the dense perturbation setting, and discriminability decreased slightly but remained above the unscaled reference value of 0.65 in the sparse case. This lack of significant decrease in discriminability across MCA perturbations suggests its utility for capturing variance within datasets without compromising the robustness and reliability of the dataset as a whole, and possibly suggests this technique as a cost effective and context-agnostic method for dataset augmentation.

While the discriminability of individuals is essential for the identification of individual brain networks, it is similarly reliant on network similarity—or lack of discriminability—across equivalent acquisitions (Hypothesis 2). In this case, connectomes were grouped based upon session, rather than subject, and the ability to distinguish one session from another based on subsamples was computed within-individual and aggregated. Both the unperturbed and dense perturbation settings perfectly preserved differences between sessions with a score of 1.0 ($p < 0.005$; optimal score: 0.5; chance: 0.5), indicating a dominant session-dependent

Table 1. The impact of instabilities as evaluated through the discriminability of the dataset based on individual (or subject) differences, session, and subsample.

Comparison	Chance	Target	Unscaled Ref.		Scaled Ref.		Dense MCA		Sparse MCA	
			Det.	Prob.	Det.	Prob.	Det.	Prob.	Det.	Prob.
H_1 : Across Subjects	0.04	1.0	0.64	0.65	0.82	0.82	0.82	0.82	0.77	0.75
H_2 : Across Sessions	0.5	0.5	1.00	1.00	1.00	1.00	1.00	1.00	0.88	0.85
H_3 : Across Subsamples	0.5	0.5					0.99	1.00	0.71	0.61

The performance is reported as mean discriminability. While a perfectly discriminable dataset would be represented by a score of 1.0, the chance performance, indicating minimal discriminability, is 1/the number of classes. H_3 could not be tested using the reference executions due to too few possible comparisons. The alternative hypothesis, indicating significant discrimination, was accepted for all experiments, with $p < 0.005$ after correcting for multiple comparisons.

<https://doi.org/10.1371/journal.pone.0250755.t001>

signal for all individuals despite no intended biological differences. However, while still significant relative to chance (score: 0.85 and 0.88; $p < 0.005$ for both), sparse perturbations lead to significantly lower discriminability of the dataset ($p < 0.005$ for all). This reduction of the difference between sessions suggests that the added variance due to perturbations reduces the relative impact of non-biological acquisition-dependent bias inherent in the networks.

Though the previous sets of experiments inextricably evaluate the interaction between data acquisition and tool, the use of subsampling allowed for characterizing the discriminability of networks sampled from within a single acquisition (Hypothesis 3). While this experiment could not be evaluated using reference executions, the networks generated with dense perturbations showed near perfect discrimination between subsamples, with scores of 0.99 and 1.0 ($p < 0.005$; optimal: 0.5; chance: 0.5). Given that there was no variability in data acquisition, due to undesired effects such as participant motion, or preprocessing, the ability to discriminate between equivalent subsamples in this experiment may only be due to instability or bias inherent to the pipelines. The high variability introduced through sparse perturbations considerably lowered the discriminability towards chance (score: 0.71 and 0.61; $p < 0.005$ for all), further supporting this as an effective method for obtaining lower-bias estimates of individual connectivity.

Across all cases, the induced perturbations maintained the ability to discriminate networks on the basis of meaningful biological signal alongside a reduction in discriminability due to off-target signal in the sparse perturbation setting. This result appears strikingly like a manifestation of the well-known bias-variance tradeoff [30] in machine learning, a concept which observes a decrease in bias as variance is favoured by a model. In particular, this highlights that numerical perturbations can be used to not only evaluate the stability of pipelines, but that the induced variance may be leveraged for the interpretation as a robust distribution of possible results.

Distributions of graph statistics are reliable, but individual statistics are not

Exploring the stability of topological features of structural connectomes is relevant for typical analyses, as low dimensional features are often more suitable than full connectomes for many analytical methods in practice [5]. A separate subset of the NKIRS dataset was randomly selected to contain a single non-sampled session for 100 individuals ($100 \times 1 \times 1$) using the pipelines and instrumentation methods to generate connectomes as above. Connectomes were generated 20 times each, resulting in a dataset which also contained 8,400 connectomes with the MCA simulations serving as the only source of repeated measurements.

The stability of several commonly-used multivariate graph features [31] were explored and are presented in Fig 2. The cumulative density of the features was computed within individuals and the mean cumulative density and associated standard error were computed for across individuals (Fig 2A and 2B). There was no significant difference between the distributions for each feature across the two perturbation settings, suggesting that the topological features summarized by these multivariate features are robust across both perturbation modes.

In addition to the comparison of distributions, the stability of the first 5 moments of these features was evaluated (Fig 2C and 2D). In the face of dense perturbations, the feature-moments were stable with more than 10 significant digits with the exception of edge weight when using the deterministic pipeline, though the probabilistic pipeline was more stable for all comparisons ($p < 0.0001$; exploratory). In stark contrast, sparse perturbations led to highly unstable feature-moments (Fig 2D), such that none contained more than 5 significant digits of information and several contained less than a single significant digit, indicating a complete

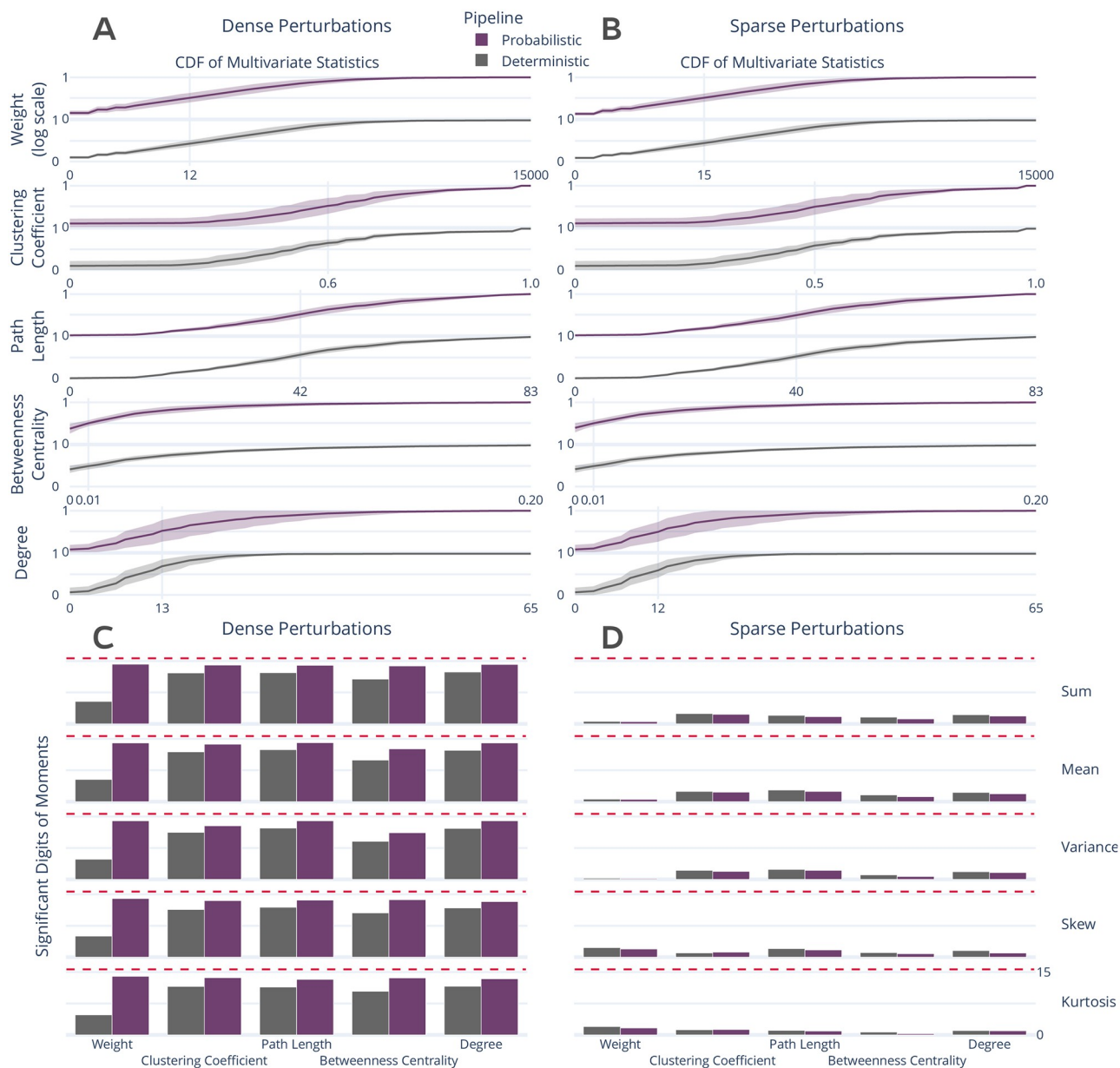


Fig 2. Distribution and stability assessment of multivariate graph statistics. (A, B) The cumulative distribution functions of multivariate statistics across all subjects and perturbation settings. There was no significant difference between the distributions in A and B. (C, D) The number of significant digits in the first five moments of each statistic across perturbations. The dashed red line refers to the maximum possible number of significant digits.

<https://doi.org/10.1371/journal.pone.0250755.g002>

lack of reliability. This dramatic degradation in stability for individual measures strongly suggests that these features may be unreliable as individual biomarkers when derived from a single pipeline evaluation, though their reliability may be increased when studying their distributions across perturbations. A similar analysis was performed for univariate statistics which obtained similar findings and can be found in S3 Section in [S1 File](#).

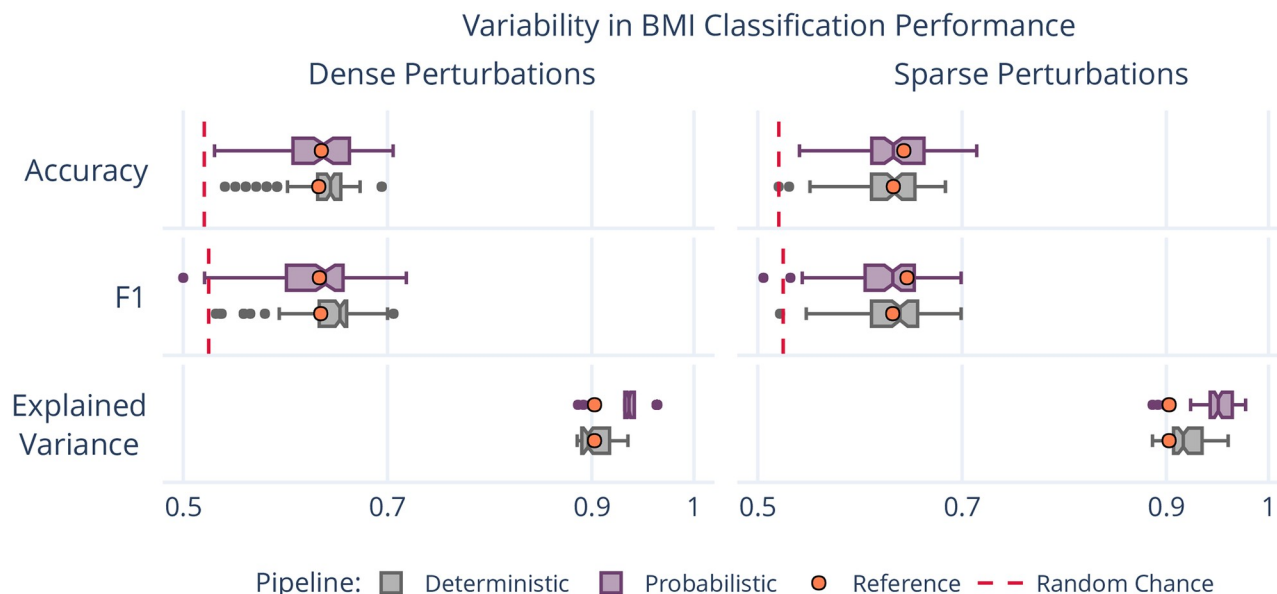


Fig 3. Variability in BMI classification across the sampling of an MCA-perturbed dataset. The dashed red lines indicate random-chance performance, and the orange dots show the performance using the reference executions.

<https://doi.org/10.1371/journal.pone.0250755.g003>

Uncertainty in brain-phenotype relationships

While the variability of connectomes and their features was summarized above, networks are commonly used as inputs to machine learning models tasked with learning brain-phenotype relationships [6]. To explore the stability of these analyses, we modelled the relationship between high- or low- Body Mass Index (BMI) groups and brain connectivity using standard dimensionality reduction and classification tools. In particular, we used Principal Component Analysis followed by a Logistic Regression classifier to predict BMI label, and demonstrated similar performance to previous work which adopted similar techniques for this task [32, 33]. We compared the performance achieved across numerically perturbed samples to both the reference and random performance (Fig 3).

The analysis was perturbed through distinct samplings of the dataset across both pipelines and perturbation methods. The accuracy and F1 score for the perturbed models varied from 0.520–0.716 and 0.510–0.725, respectively, ranging from at or below random performance to outperforming performance on the reference dataset. This large variability illustrates a previously uncharacterized margin of uncertainty in the modelling of this relationship, and limits confidence in reported accuracy scores on singly processed datasets. The portion of explained variance in these samples ranged from 88.6% – 97.8%, similar to the reference of 90.3%, suggesting that the range in performance was not due to a gain or loss of meaningful signal, but rather the reduction of bias towards specific outcome. Importantly, this finding does not suggest that modelling brain-phenotype relationships is not possible, but rather it sheds light on impactful uncertainty that must be accounted for in this process, and supports the use of ensemble modeling techniques.

One distinction between the results presented here and the previous is that while networks derived from dense perturbations had been shown to exhibit less dramatic instabilities in general, the results here show similar variability in classification performance across the two methods. This consistency suggests that the desired method of instrumentation may vary across

experiments. While sparse perturbations result in considerably more variability in networks directly, the two techniques capture similar variability when relating networks to this phenotypic variable. Given the dramatic reduction in computational overhead, a sparse instrumentation may be preferred when processing datasets for eventual application in modelling brain-phenotype relationships.

Discussion

The perturbation of structural connectome estimation pipelines with small amounts of noise, on the order of machine error, led to considerable variability in derived brain graphs. Across all analyses the stability of results ranged from nearly perfectly trustworthy (i.e. no variation) to completely unreliable (i.e. containing no trustworthy information). Given that the magnitude of introduced numerical noise is to be expected in computational workflows, this finding has potentially significant implications for inferences in brain imaging as it is currently performed. In particular, this bounds the success of studying individual differences, a central objective in brain imaging [6], given that the quality of relationships between phenotypic data and brain networks will be limited by the stability of the connectomes themselves. This issue is accentuated through the crucial finding that individually derived network features were unreliable despite there being no significant difference in their aggregated distributions. This finding is not damning for the study of brain networks as a whole, but rather is strong support for the aggregation of networks, either across perturbations for an individual or across groups, over the use of individual estimates.

Underestimated false positive rates

While the instability of brain networks was used here to demonstrate the limitations of modelling brain-phenotype relationships in the context of machine learning, this limitation extends to classical hypothesis testing, as well. Though performing individual comparisons in a hypothesis testing framework will be accompanied by reported false positive rates, the accuracy of these rates is critically dependent upon the reliability of the samples used. In reality, the true false positive rate for a test would be a combination of the reported confidence and the underlying variability in the results, a typically unknown quantity.

When performing these experiments outside of a repeated-measure context, such as that afforded here through MCA, it is impossible to empirically estimate the reliability of samples. This means that the reliability of accepted hypotheses is also unknown, regardless of the reported false positive rate. In fact, it is a virtual certainty that the true false positive rate for a given hypothesis exceeds the reported value simply as a result of numerical instabilities. This uncertainty inherent to derived data is compounded with traditional arguments limiting the trustworthiness of claims [34], and hampers the ability of researchers to evaluate the quality of results. The accompaniment of brain imaging experiments with direct evaluations of their stability, as was done here, would allow researchers to simultaneously improve the numerical stability of their analyses and accurately gauge confidence in them. The induced variability in derived brain networks may be leveraged to estimate aggregate connectomes with lower bias than any single independent observation, leading to learned relationships that are more generalizable and ultimately more useful.

Cost-effective data augmentation

The evaluation of reliability in brain imaging has historically relied upon the expensive collection of repeated measurements choreographed by massive cross-institutional consortia [35, 36]. The finding that perturbing experiments using MCA both preserved the discriminability

of the dataset due to biological signal and decreased the discriminability due to off-target differences across acquisitions and subsamples opens the door for a promising paradigm shift. Given that MCA is data-agnostic, this technique could be used effectively in conjunction with, or in lieu of, realistic noise models to augment existing datasets. While this of course would not replace the need for repeated measurements when exploring the effect of data collection paradigm or study longitudinal progressions of development or disease, it could be used in conjunction with these efforts to decrease the bias of each distinct sample within a dataset. In contexts where repeated measurements are typically collected to increase the fidelity of the dataset, MCA could potentially serve as an alternative solution to capture more biological variability, with the added benefit being the savings of millions of dollars on data collection.

Shortcomings and future questions

Given the complexity of recompiling complex software libraries, pre-processing was not perturbed in these experiments as the instrumentation of the canonical workflow used in diffusion image processing would have added considerable technical complexity and computational overhead to the large set of experiments performed here; similarly, this complexity along with the added layer of difficulty in comparing instrumentations meant that only algorithms within a single library were tested. Other work has shown that linear registration, a core piece of many elements of pre-processing such as motion correction and alignment, is sensitive to minor perturbations [21]. It is likely that the instabilities across the entire processing workflow would be compounded with one another, resulting in even greater variability. While the analyses performed in this paper evaluated a single dataset and set of pipelines, extending this work to other modalities and analyses, alongside the detection of local sources of instability within pipelines, is of interest for future projects.

This paper does not explore methodological flexibility or compare this to numerical instability. Recently, the nearly boundless space of analysis pipelines and their impact on outcomes in brain imaging has been clearly demonstrated [17]. The approach taken in these studies complement one another and explore instability at the opposite ends of the spectrum, with human variability in the construction of an analysis workflow on one end and the unavoidable error implicit in the digital representation of data on the other. It is of extreme interest to combine these approaches and explore the interaction of these scientific degrees of freedom with effects from software implementations, libraries, and parametric choices.

Finally, it is important to state explicitly that the work presented here does not invalidate analytical pipelines used in brain imaging, but merely sheds light on the fact that many studies are accompanied by an unknown degree of uncertainty due to machine-introduced errors. The presence of unknown error-bars associated with experimental findings limits the impact of results due to increased uncertainty. The desired outcome of this paper is to motivate a shift in scientific computing—both in neuroimaging and more broadly—towards a paradigm that favours the explicit evaluation of the trustworthiness of claims alongside the claims themselves.

Materials & methods

Dataset

The Nathan Kline Institute Rockland Sample (NKI-RS) [27] dataset contains high-fidelity imaging and phenotypic data from over 1,000 individuals spread across the lifespan. A subset of this dataset was chosen for each experiment to both match sample sizes presented in the original analyses and to minimize the computational burden of performing MCA. The selected

subset comprises 100 individuals ranging in age from 6–79 with a mean of 36.8 (original: 6–81, mean 37.8), 60% female (original: 60%), with 52% having a BMI over 25 (original: 54%).

Each selected individual had at least a single session of both structural T1-weighted (MPRAGE) and diffusion-weighted (DWI) MR imaging data. DWI data was acquired with 137 diffusion directions in a single shell; more information regarding the acquisition of this dataset can be found in the NKI-RS data release [27].

In addition to the 100 sessions mentioned above, 25 individuals had a second session to be used in a test-retest analysis. Two additional copies of the data for these individuals were generated, including only the odd or even diffusion directions ($64 + 9 \text{ B0 volumes} = 73$ in either case) such that the acquired data was evenly represented across both portions, and each subsample represented a realistic complete acquisition. This allowed for an extra level of stability evaluation to be performed between the levels of MCA and session-level variation.

In total, the dataset is composed of 100 subsampled sessions of data originating from 50 acquisitions and 25 individuals for in depth stability analysis, and an additional 100 sessions of full-resolution data from 100 individuals for subsequent analyses.

Processing

The dataset was preprocessed using a standard FSL [37] workflow consisting of eddy-current correction and alignment. The MNI152 atlas [38] was aligned to each session of data via the structural images, and the resulting transformation was applied to the DKT parcellation [39]. Subsampling the diffusion data took place after preprocessing was performed on full-resolution sessions, ensuring that an additional confound was not introduced in this process when comparing between downsampled sessions. The preprocessing described here was performed once without MCA, and thus is not being evaluated.

Structural connectomes were generated from preprocessed data using two canonical pipelines from Dipy [28]: deterministic and probabilistic. In the deterministic pipeline, a constant solid angle model was used to estimate tensors at each voxel and streamlines were then generated using the EuDX algorithm [29]. In the probabilistic pipeline, a constrained spherical deconvolution model was fit at each voxel and streamlines were generated by iteratively sampling the resulting fiber orientation distributions. In both cases tracking occurred with 8 seeds per 3D voxel and edges were added to the graph based on the location of terminal nodes with weight determined by fiber count.

The random state of both pipelines was fixed for all analyses. Fixing this random state led to entirely deterministic repeated-evaluations of the tools, and allowed for explicit attribution of observed variability to limitations in tool precision as provoked by Monte Carlo simulations, rather than the internal state of the algorithm.

Perturbations

All connectomes were generated with one reference execution where no perturbation was introduced in the processing. For all other executions, all floating point operations were instrumented with Monte Carlo Arithmetic (MCA) [25] through Verificarlo [26]. MCA simulates the distribution of errors implicit to all instrumented floating point operations (flop). This rounding is performed on a value x at precision t by:

$$\text{inexact}(x) = x + 2^{e_x - t} \xi \quad (1)$$

where e_x is the exponent value of x and ξ is a uniform random variable in the range $(-\frac{1}{2}, \frac{1}{2})$.

MCA can be introduced in two places for each flop: before or after evaluation. Performing MCA on the inputs of an operation limits its precision, while performing MCA on the output

of an operation highlights round-off errors that may be introduced. The former is referred to as Precision Bounding (PB) and the latter is called Random Rounding (RR).

Using MCA, the execution of a pipeline may be performed many times to produce a distribution of results. Studying the distribution of these results can then lead to insights on the stability of the instrumented tools or functions. To this end, a complete software stack was instrumented with MCA and is made available on GitHub at <https://github.com/verificarlo/fuzzy>.

The RR variant of MCA was used for all experiments. As was presented in [19], both the degree of instrumentation (i.e. number of affected libraries) and the perturbation mode have an effect on the distribution of observed results. For this work, the RR-MCA was applied across the bulk of the relevant operations (those occurring in BLAS, LAPACK, Python, Cython, and Numpy) and is referred to as dense perturbation. In this case the bulk of numerical operations were affected by MCA.

Conversely, the case in which RR-MCA was applied across the operations in a small subset of operations (those occurring in Python and Cython) is here referred to as sparse perturbation. In this case, the inputs to operations within the instrumented libraries were perturbed, resulting in less frequent, data-centric perturbations. Alongside the stated theoretical differences, sparse perturbation is considerably less computationally expensive than dense perturbation.

All perturbations targeted the least-significant-bit for all data ($t = 24$ and $t = 53$ in float32 and float64, respectively [26]). Perturbing the least significant bit importantly serves as a perturbation of machine error, and thus is the appropriate precision to be applied globally in complex pipelines. Simulations were performed 20 times for each pipeline execution for the 100 sample dataset and 10 times for the repeated measures dataset. A detailed motivation for the number of simulations can be found in [40].

Evaluation

The magnitude and importance of instabilities in pipelines can be considered at a number of analytical levels, namely: the induced variability of derivatives directly, the resulting downstream impact on summary statistics or features, or the ultimate change in analyses or findings. We explore the nature and severity of instabilities through each of these lenses. Unless otherwise stated, all p-values were computed using Wilcoxon signed-rank tests and corrected for multiple comparisons. To avoid biasing these statistics in this unique repeated-measures context, tests were performed across sets of independent observations and then the results were aggregated in all cases.

Direct evaluation of the graphs. The differences between perturbation-generated graphs was measured directly through both a direct variance quantification and a comparison to other sources of variance such as individual- and session-level differences.

Quantification of variability. Graphs, in the form of adjacency matrices, were compared to one another using three metrics: normalized percent deviation, Pearson correlation, and edge-wise significant digits. The normalized percent deviation measure, defined in [19], scales the norm of the difference between a simulated graph and the reference execution (that without intentional perturbation) with respect to the norm of the reference graph, and is defined as [19]:

$$\%Dev(A, B) = \sqrt{\sum_{i=1}^m \sum_{j=1}^n |a_{ij} - b_{ij}|^2} / \sqrt{\sum_{i=1}^m \sum_{j=1}^n |a_{ij}|^2}, \quad (2)$$

where A and B each represent a graph, and a_{ij} are elements therein corresponding to row and column i and j , respectively. For these experiments, the A graph always refers to the reference,

where B represents a perturbed value. The purpose of this comparison is to provide insight on the scale of differences in observed graphs relative to the original signal intensity. A Pearson correlation coefficient [41] was computed in complement to normalized percent deviation to identify the consistency of structure and not just intensity between observed graphs, though the result of this experiment is shown only in S1 Section in [S1 File](#).

Finally, the estimated number of significant digits, s' , for each edge in the graph is calculated as:

$$s' = -\log_{10} \frac{\sigma}{|\mu|} \quad (3)$$

where μ and σ are the mean and unbiased estimator of standard deviation across graphs, respectively. The upper bound on significant digits is 15.7 for 64-bit floating point data.

The percent deviation, correlation, and number of significant digits were each calculated within a single session of data, thereby removing any subject- and session-effects and providing a direct measure of the tool-introduced variability across perturbations. A distribution was formed by aggregating these individual results.

Class-based variability evaluation. To gain a concrete understanding of the significance of observed variations we explore the separability of our results with respect to understood sources of variability, such as subject-, session-, and pipeline-level effects. This can be probed through Discriminability [14], a technique similar to ICC [12] which relies on the mean of a ranked distribution of distances between observations belonging to a defined set of classes. The discriminability statistic is formalized as follows:

$$Disc. = Pr(\|g_{ij} - g_{i'j'}\| \leq \|g_{ij} - g_{i'j'}\|) \quad (4)$$

where g_{ij} is a graph belonging to class i that was measured at observation j , where $i \neq i'$ and $j \neq j'$.

Discriminability can then be read as the probability that an observation belonging to a given class will be more similar to other observations within that class than observations of a different class. It is a measure of reproducibility, and is discussed in detail in [14]. This definition allows for the exploration of deviations across arbitrarily defined classes that in practice can be any of those listed above. We combine this statistic with permutation testing to test hypotheses on whether differences between classes are statistically significant in each of these settings. This statistic is similar to ICC [12] in a two-measurement setting, however, given the dependence on a rank distribution from all measurements, discriminability scores do not become meaningless by the addition of more samples which are highly similar to the originals, whereas ICC scores would much more rapidly trend towards 1, making discriminability appropriate in this context. The scaling properties of discriminability are described more fully in S2 Section in [S1 File](#).

With this in mind, three hypotheses were defined. For each setting, we state the alternate hypotheses, the variable(s) which were used to determine class membership, and the remaining variables which may be sampled when obtaining multiple observations. Each hypothesis was tested independently for each pipeline and perturbation mode.

H_{A1} : Individuals are distinct from one another

Class definition: *Subject ID*

Comparisons: **Session (1 subsample)**, *Subsample (1 session)*, *MCA (1 subsample, 1 session)*

H_{A2} : Sessions within an individual are distinct

Class definition: *Session ID | Subject ID*

Comparisons: **Subsample**, *MCA (1 subsample)*

*H*_{A3}: Subsamples are distinctClass definition: *Subsample* | *Subject ID*, *Session ID*Comparisons: **MCA**

As a result, we tested 3 hypotheses across 6 MCA experiments and 3 reference experiments on 2 pipelines and 2 perturbation modes, resulting in a total of 30 distinct tests. While results from all tests can be found within S2 Section in [S1 File](#), only the bolded comparisons in the list above have been presented in the main body of this article. Correction for repeated testing was performed.

Evaluating graph-theoretical metrics. While connectomes may be used directly for some analyses, it is common practice to summarize them with structural measures, that can then be used as lower-dimensional proxies of connectivity in so-called graph-theoretical studies [5]. We explored the stability of several commonly-used univariate (graphwise) and multivariate (nodewise or edgewise) features. The features computed and subsequent methods for comparison in this section were selected to closely match those computed in [31].

Univariate differences. For each univariate statistic (edge count, mean clustering coefficient, global efficiency, modularity of the largest connected component, assortativity, and mean path length) a distribution of values across all perturbations within subjects was observed. A Z-score was computed for each sample with respect to the distribution of feature values within an individual, and the proportion of “classically significant” Z-scores, i.e. corresponding to $p < 0.05$, was reported and aggregated across all subjects. There was no correction for multiple comparisons in these statistics, as they were not used to interpret a hypothesis but demonstrate the false-positive rate due to perturbations. The number of significant digits contained within an estimate derived from a single subject were calculated and aggregated. The results of this analysis can be found in S3 Section in [S1 File](#).

Multivariate differences. In the case of both nodewise (degree distribution, clustering coefficient, betweenness centrality) and edgewise (weight distribution, connection length) features, the cumulative density functions of their distributions were evaluated over a fixed range and subsequently aggregated across individuals. The number of significant digits for each moment of these distributions (sum, mean, variance, skew, and kurtosis) were calculated across observations within a sample and aggregated.

Evaluating a brain-phenotype analysis. Though each of the above approaches explores the instability of derived connectomes and their features, many modern studies employ modeling or machine-learning approaches, for instance to learn brain-phenotype relationships or identify differences across groups. We carried out one such study and explored the instability of its results with respect to the upstream variability of connectomes characterized in the previous sections. We performed the modeling task with a single sampled connectome per individual and repeated this sampling and modelling 20 times. We report the model performance for each sampling of the dataset and summarize its variance.

BMI classification. Structural changes have been linked to obesity in adolescents and adults [42]. We classified normal-weight and overweight individuals from their structural networks (using for overweight a cutoff of BMI >25 [33]). We reduced the dimensionality of the connectomes through principal component analysis (PCA), and provided the first N-components to a logistic regression classifier for predicting BMI class membership, similar to methods shown in [32, 33]. The number of components was selected as the minimum set which explained >90% of the variance when averaged across the training set for each fold within the cross validation of the original graphs; this resulted in a feature of 20 components. We trained the model using *k*-fold cross validation, with $k = 2, 5, 10$, and *N* (equivalent to leave-one-out; LOO).

Data & code provenance. The unprocessed dataset is available through The Consortium of Reliability and Reproducibility (http://fcon_1000.projects.nitrc.org/indi/enhanced/), including both the imaging data as well as phenotypic data which may be obtained upon submission and compliance with a Data Usage Agreement. The connectomes generated through simulations have been bundled and stored permanently (<https://doi.org/10.5281/zenodo.4041549>), and are made available through The Canadian Open Neuroscience Platform (<https://portal.conp.ca/search>, search term “Kiar”).

All software developed for processing or evaluation is publicly available on GitHub at <https://github.com/gkpapers/2021ImpactOfInstability>. Experiments were launched using Boutiques [43] and Clowdr [44] in Compute Canada’s HPC cluster environment. MCA instrumentation was achieved through Verificarlo [26] available on Github at <https://github.com/verificarlo/verificarlo>. A set of MCA instrumented software containers is available on Github at <https://github.com/gkiar/fuzzy>.

Supporting information

S1 File.
(PDF)

Acknowledgments

This work was supported in partnership with Health Canada, for the Canadian Open Neuroscience Platform initiative.

Author Contributions

Conceptualization: Gregory Kiar, Alan C. Evans, Tristan Glatard.

Data curation: Gregory Kiar.

Formal analysis: Gregory Kiar.

Funding acquisition: Gregory Kiar.

Investigation: Gregory Kiar, Tristan Glatard.

Methodology: Gregory Kiar, Yohan Chatelain, Pablo de Oliveira Castro, Eric Petit, Ariel Rokem, Gaël Varoquaux, Bratislav Misic, Tristan Glatard.

Project administration: Gregory Kiar.

Resources: Gregory Kiar.

Software: Gregory Kiar, Yohan Chatelain, Pablo de Oliveira Castro, Eric Petit.

Supervision: Alan C. Evans, Tristan Glatard.

Validation: Gregory Kiar, Pablo de Oliveira Castro, Eric Petit.

Visualization: Gregory Kiar.

Writing – original draft: Gregory Kiar.

Writing – review & editing: Gregory Kiar, Ariel Rokem, Gaël Varoquaux, Bratislav Misic, Alan C. Evans, Tristan Glatard.

References

- Behrens T. E. and Sporns O., "Human connectomics," *Current opinion in neurobiology*, vol. 22, no. 1, pp. 144–153, 2012. <https://doi.org/10.1016/j.conb.2011.08.005> PMID: 21908183
- Xia M., Lin Q., Bi Y., and He Y., "Connectomic insights into topologically centralized network edges and relevant motifs in the human brain," *Frontiers in human neuroscience*, vol. 10, p. 158, 2016. <https://doi.org/10.3389/fnhum.2016.00158>
- Morgan J. L. and Lichtman J. W., "Why not connectomics?" *Nature methods*, vol. 10, no. 6, p. 494, 2013. <https://doi.org/10.1038/nmeth.2480>
- Van den Heuvel M. P., Bullmore E. T., and Sporns O., "Comparative connectomics," *Trends in cognitive sciences*, vol. 20, no. 5, pp. 345–361, 2016. <https://doi.org/10.1016/j.tics.2016.03.001>
- Rubinov M. and Sporns O., "Complex network measures of brain connectivity: uses and interpretations," *Neuroimage*, vol. 52, no. 3, pp. 1059–1069, Sep. 2010. <https://doi.org/10.1016/j.neuroimage.2009.10.003>
- Dubois J. and Adolphs R., "Building a science of individual differences from fMRI," *Trends Cogn. Sci.*, vol. 20, no. 6, pp. 425–443, Jun. 2016. <https://doi.org/10.1016/j.tics.2016.03.014>
- Fornito A. and E. T. Bullmore, "Connectomics: a new paradigm for understanding brain disease," *European Neuropsychopharmacology*, vol. 25, no. 5, pp. 733–748, 2015. <https://doi.org/10.1016/j.euroneuro.2014.02.011>
- Deco G. and Kringelbach M. L., "Great expectations: using whole-brain computational connectomics for understanding neuropsychiatric disorders," *Neuron*, vol. 84, no. 5, pp. 892–905, 2014. <https://doi.org/10.1016/j.neuron.2014.08.034>
- Xie T. and He Y., "Mapping the alzheimer's brain with connectomics," *Frontiers in psychiatry*, vol. 2, p. 77, 2012. <https://doi.org/10.3389/fpsy.2011.00077>
- Filippi M., van den Heuvel M. P., Fornito A., He Y., Pol H. E. H., Agosta F., et al., "Assessment of system dysfunction in the brain through mri-based connectomics," *The Lancet Neurology*, vol. 12, no. 12, pp. 1189–1199, 2013. [https://doi.org/10.1016/S1474-4422\(13\)70144-3](https://doi.org/10.1016/S1474-4422(13)70144-3)
- Van Den Heuvel M. P. and Fornito A., "Brain networks in schizophrenia," *Neuropsychology review*, vol. 24, no. 1, pp. 32–48, 2014. <https://doi.org/10.1007/s11065-014-9248-7>
- Bartko J. J., "The intraclass correlation coefficient as a measure of reliability," *Psychol. Rep.*, vol. 19, no. 1, pp. 3–11, Aug. 1966. <https://doi.org/10.2466/pr0.1966.19.1.3> PMID: 5942109
- Brandmaier A. M., Wenger E., Bodammer N. C., Kühn S., Raz N., and Lindenberger U., "Assessing reliability in neuroimaging research through intra-class effect decomposition (ICED)," *Elife*, vol. 7, Jul. 2018. <https://doi.org/10.7554/eLife.35718>
- E. W. Bridgeford, S. Wang, Z. Yang, Z. Wang, T. Xu, C. Craddock, et al., "Eliminating accidental deviations to minimize generalization error: applications in connectomics and genomics," *bioRxiv*, p. 802629, 2020.
- G. Kiar, E. Bridgeford, W. G. Roncal, V. Chandrashekar, and others, "A High-Throughput pipeline identifies robust connectomes but troublesome variability," *bioRxiv*, 2018.
- Antonakakis M., Schrader S., Aydın Ü., Khan A., Gross J., Zervakis M., et al., "Inter-subject variability of skull conductivity and thickness in calibrated realistic head models," *Neuroimage*, vol. 223, p. 117353, 2020. <https://doi.org/10.1016/j.neuroimage.2020.117353>
- Botvinik-Nezer R., Holzmeister F., Camerer C. F., Dreber A., Huber J., Johannesson M., et al., "Variability in the analysis of a single neuroimaging dataset by many teams," *Nature*, pp. 1–7, 2020.
- Eklund A., Nichols T. E., et al., "Cluster failure: Why fMRI inferences for spatial extent have inflated false-positive rates," *Proceedings of the national academy of sciences*, vol. 113, no. 28, pp. 7900–7905, 2016. <https://doi.org/10.1073/pnas.1602413113>
- G. Kiar, P. de Oliveira Castro, P. Rioux, E. Petit, S. T. Brown, A. C. Evans, et al., "Comparing perturbation models for evaluating stability of neuroimaging pipelines," *The International Journal of High Performance Computing Applications*, 2020.
- L. B. Lewis, C. Y. Lepage, N. Khalili-Mahani, M. Omidyeganeh, S. Jeon, P. Bermudez, et al., "Robustness and reliability of cortical surface reconstruction in CIVET and FreeSurfer," *Annual Meeting of the Organization for Human Brain Mapping*, 2017.
- Glatard T., Lewis L. B., Ferreira da Silva R., Adalat R., Beck N., Lepage C., et al., "Reproducibility of neuroimaging analyses across operating systems," *Front. Neuroinform.*, vol. 9, p. 12, Apr. 2015. <https://doi.org/10.3389/fninf.2015.00012>
- A. Salari, G. Kiar, L. Lewis, A. C. Evans, and T. Glatard, "File-based localization of numerical perturbations in data analysis pipelines," *arXiv preprint arXiv:2006.04684*, 2020.

23. Bennett C. M., Miller M. B., Wolford G. L., "Neural correlates of interspecies perspective taking in the post-mortem Atlantic Salmon: An argument for multiple comparisons correction," *Neuroimage*, vol. 47, pp. S125, 2009. [https://doi.org/10.1016/S1053-8119\(09\)71202-9](https://doi.org/10.1016/S1053-8119(09)71202-9)
24. M. Baker, "1,500 scientists lift the lid on reproducibility," *Nature*, 2016.
25. D. S. Parker, *Monte Carlo Arithmetic: exploiting randomness in floating-point arithmetic*. University of California (Los Angeles). Computer Science Department, 1997.
26. C. Denis, P. de Oliveira Castro, and E. Petit, "Verificarlo: Checking floating point accuracy through monte carlo arithmetic," *2016 IEEE 23rd Symposium on Computer Arithmetic (ARITH)*, 2016.
27. Nooner K. B., Colcombe S. J., Tobe R. H., Mennes M. et al., "The NKI-Rockland sample: A model for accelerating the pace of discovery science in psychiatry," *Front. Neurosci.*, vol. 6, p. 152, Oct. 2012. <https://doi.org/10.3389/fnins.2012.00152>
28. Garyfallidis E., Brett M., Amirkhian B., Rokem A., van der Walt S., Descoteaux M., et al, "Dipy, a library for the analysis of diffusion MRI data," *Front. Neuroinform.*, vol. 8, p. 8, Feb. 2014. <https://doi.org/10.3389/fninf.2014.00008>
29. Garyfallidis E., Brett M., Correia M. M., Williams G. B., and Nimmo-Smith I., "QuickBundles, a method for tractography simplification," *Front. Neurosci.*, vol. 6, p. 175, Dec. 2012. <https://doi.org/10.3389/fnins.2012.00175>
30. Geman S., Bienenstock E., and Doursat R., "Neural networks and the bias/variance dilemma," *Neural computation*, vol. 4, no. 1, pp. 1–58, 1992. <https://doi.org/10.1162/neco.1992.4.1.1>
31. Betzel R. F., Griffa A., Hagmann P., and Mišić B., "Distance-dependent consensus thresholds for generating group-representative structural brain networks," *Network neuroscience*, vol. 3, no. 2, pp. 475–496, 2019. https://doi.org/10.1162/netn_a_00075 PMID: 30984903
32. Park B.-Y., Seo J., Yi J., and Park H., "Structural and functional brain connectivity of people with obesity and prediction of body mass index using connectivity," *PLoS One*, vol. 10, no. 11, p. e0141376, Nov. 2015. <https://doi.org/10.1371/journal.pone.0141376>
33. Gupta A., Mayer E. A., Sanmiguel C. P., Van Horn J. D., Woodworth D., Ellingson B. M., et al, "Patterns of brain structural connectivity differentiate normal weight from overweight subjects," *Neuroimage Clin*, vol. 7, pp. 506–517, Jan. 2015. <https://doi.org/10.1016/j.nicl.2015.01.005>
34. Ioannidis J. P., "Why most published research findings are false," *PLoS medicine*, vol. 2, no. 8, p. e124, 2005. <https://doi.org/10.1371/journal.pmed.0020124> PMID: 16060722
35. Van Essen D. C., Smith S. M., Barch D. M., Behrens T. E., Yacoub E., Ugurbil K., et al., "The WU-Minn human connectome project: an overview," *Neuroimage*, vol. 80, pp. 62–79, 2013. <https://doi.org/10.1016/j.neuroimage.2013.05.041> PMID: 23684880
36. Zuo X.-N., Anderson J. S., Bellec P., Birn R. M., Biswal B. B., Blautzik J., et al., "An open science resource for establishing reliability and reproducibility in functional connectomics," *Scientific data*, vol. 1, no. 1, pp. 1–13, 2014. <https://doi.org/10.1038/sdata.2014.49>
37. Jenkinson M., Beckmann C. F., Behrens T. E. J., Woolrich M. W., and Smith S. M., "FSL," *Neuroimage*, vol. 62, no. 2, pp. 782–790, Aug. 2012. <https://doi.org/10.1016/j.neuroimage.2011.09.015>
38. Lancaster J. L., Tordesillas-Gutiérrez D., Martinez M., Salinas F., Evans A., Zilles K., et al, "Bias between mni and talairach coordinates analyzed using the icbm-152 brain template," *Human brain mapping*, vol. 28, no. 11, pp. 1194–1205, 2007. <https://doi.org/10.1002/hbm.20345> PMID: 17266101
39. Klein A. and Tourville J., "101 labeled brain images and a consistent human cortical labeling protocol," *Front. Neurosci.*, vol. 6, p. 171, Dec. 2012. <https://doi.org/10.3389/fnins.2012.00171>
40. D. Sohler, P. De Oliveira Castro, F. Févotte, B. Lathuilière, E. Petit, and O. Jamond, "Confidence intervals for stochastic arithmetic," Jul. 2018.
41. Benesty J., Chen J., Huang Y., and Cohen I., "Pearson correlation coefficient," in *Noise Reduction in Speech Processing*, Cohen I., Huang Y., Chen J., and Benesty J., Eds. Berlin Heidelberg: Springer Berlin Heidelberg, 2009, pp. 1–4.
42. Raji C. A., Ho A. J., Parikshak N. N., Becker J. T., Lopez O. L., Kuller L. H., et al, "Brain structure and obesity," *Hum. Brain Mapp.*, vol. 31, no. 3, pp. 353–364, Mar. 2010. <https://doi.org/10.1002/hbm.20870> PMID: 19662657
43. Glatard T., Kiar G., Aumentado-Armstrong T., Beck N., Bellec P., Bernard R., et al, "Boutiques: a flexible framework to integrate command-line applications in computing platforms," *Gigascience*, vol. 7, no. 5, May 2018. <https://doi.org/10.1093/gigascience/giy016>
44. Kiar G., Brown S. T., Glatard T., and Evans A. C., "A serverless tool for platform agnostic computational experiment management," *Front. Neuroinform.*, vol. 13, p. 12, Mar. 2019. <https://doi.org/10.3389/fninf.2019.00012>