

Autoreject: Automated artifact rejection for MEG and EEG data

Mainak Jas[‡], Denis A. Engemann[‡], Yousra Bekhti¹, Federico Raimondo^{3, 4, 5, 6} and Alexandre Gramfort[‡]

¹LTCI, Télécom ParisTech, Université Paris-Saclay, France

²Cognitive Neuroimaging Unit, CEA DSV/I2BM, INSERM, Université Paris-Sud., Université Paris-Saclay, NeuroSpin center, 91191 Gif/Yvette, France

³Laboratorio de Inteligencia Artificial Aplicada, Departamento de Computación, FCEyN, Universidad de Buenos Aires, Argentina

⁴CONICET, Argentina

⁵Sorbonne Universités, UPMC Univ Paris 06, Faculté de Médecine Pitié-Salpêtrière, Paris, France

⁶Institut du Cerveau et de la Moelle épinière, ICM, PICNIC Lab, F-75013, Paris, France

Abstract

We present an automated algorithm for unified rejection and repair of bad trials in magnetoencephalography (MEG) and electroencephalography (EEG) signals. Our method capitalizes on cross-validation in conjunction with a robust evaluation metric to estimate the optimal peak-to-peak threshold – a quantity commonly used for identifying bad trials in M/EEG. This approach is then extended to a more sophisticated algorithm which estimates this threshold for each sensor yielding trial-wise bad sensors. Depending on the number of bad sensors, the trial is then repaired by interpolation or by excluding it from subsequent analysis. All steps of the algorithm are fully automated thus lending itself to the name *Autoreject*.

In order to assess the practical significance of the algorithm, we conducted extensive validation and comparison with state-of-the-art methods on four public datasets containing MEG and EEG recordings from more than 200 subjects. Comparison include purely qualitative efforts as well as quantitatively benchmarking against human supervised and semi-automated preprocessing pipelines. The algorithm allowed us to automate the preprocessing of MEG data from the Human Connectome Project (HCP) going up to the computation of the evoked responses. The automated nature of our method minimizes the burden of human inspection, hence supporting scalability and reliability demanded by data analysis in modern neuroscience.

Keywords: Magnetoencephalography (MEG), Electroencephalogram (EEG), preprocessing, statistical learning, cross-validation, automated analysis, Human Connectome Project (HCP)

1 Introduction

Magneto-/electroencephalography (M/EEG) offer the unique ability to explore and study, non-invasively, the temporal dynamics of the brain and its cognitive processes. The M/EEG community has only recently begun to appreciate the importance of large-scale studies, in an effort to improve replicability and statistical power of

*46 Rue Barrault, Télécom ParisTech, Université Paris-Saclay, France. E-mail address: mainak.jas@telecom-paristech.fr

[‡]These authors jointly supervised the work

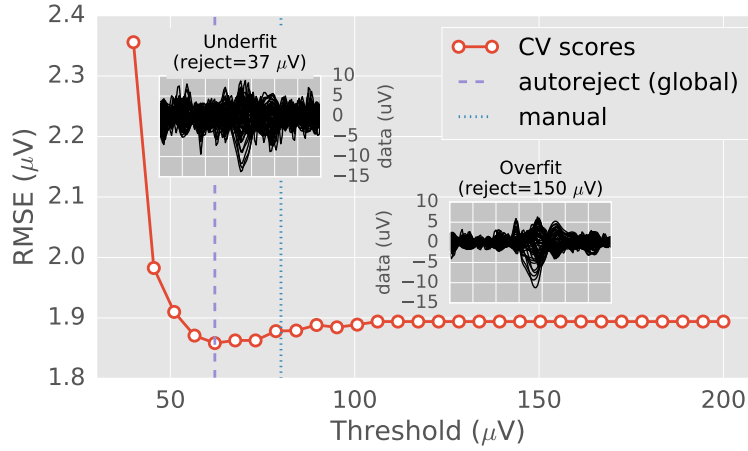


Fig. 1: Cross-validation error as a function of peak-to-peak rejection threshold on one EEG dataset. The root mean squared error (RMSE) between the mean of the training set (after removing the trials marked as bad) and the median of the validation set was used as the cross-validation metric (Section 2.1). The two insets show the average of the trials as “butterfly plots” (each curve representing one sensor) for very low and high thresholds. For low thresholds, the RMSE is high because most of the trials are rejected (underfit). At high thresholds, the model does not drop any trials (overfit). The optimal data-driven threshold (*autoreject, global*) with minimum RMSE is somewhere in between. It closely matches the human threshold.

experiments. This has given rise to the practice of sharing and publishing data in open archives (Gorgolewski and Poldrack, 2016). Examples of such large electrophysiological datasets include the Human Connectome Project (HCP) (Van Essen et al., 2012; Larson-Prior et al., 2013), the Physiobank (Goldberger et al., 2000) and the OMEGA archive (Niso et al., 2016). A tendency towards ever-growing massive datasets as well as a shift towards common standards for accessing these databases (Gorgolewski et al., 2016; Bigdely-Shamlo et al., 2013) is clearly visible. The UK Biobank project (Ollier et al., 2005) which currently hosts data from more than 50,000 subjects is yet another example of this trend.

This has however, given rise to new challenges including automating the analysis pipeline (Gorgolewski and Poldrack, 2016). Automation will not only save time, but also allow scalable analysis and reduce the barriers to reanalysis of data, thus facilitating reproducibility. Engemann and Gramfort (2015) have recently worked towards more automation in M/EEG analysis pipelines by considering the problem of covariance estimation, a step commonly done prior to source localization. Yet, one of the most critical bottlenecks that limits the reanalysis of M/EEG data remains at the preprocessing stage with the annotation and rejection of artifacts. Despite being so fundamental to M/EEG analysis given how easily such data can be corrupted by noise and artifacts, there is currently no consensus in the community on how to address this particular issue.

In the presence of what we will refer to as *bad* data, various data cleaning strategies have been employed. A first intuitive strategy is to exclude bad data from analysis, to *reject* it. While this approach is very often employed, for example, because data cleaning is time consuming, or out of reach for practitioners, it leads to a loss of data that are costly to acquire. This is particularly the case for clinical studies, where patients have difficulties staying still or focusing on the task (Cruse et al., 2012; Goldfine et al., 2013), or even when babies are involved as subjects (Basirat et al., 2014).

When working with M/EEG, the data can be bad due to the presence of bad sensors (also known as channels) and bad trials. A trial refers here to a data segment whose location in time is typically related to an experimental protocol. But here we will also call trial any data segment even if it is acquired during a task-free protocol. Accordingly, a bad trial or bad sensor is one which contains bad data. Ignoring the presence of bad data can adversely affect analysis downstream in the pipeline. For example, when multiple trials time-locked to the stimulation are averaged to estimate an evoked response, ignoring the presence of a single bad trial can corrupt the average. The mean of a random vector is not robust to the presence of strong outliers. Another example quite common in practice, both in the case of EEG and MEG, is the presence of a bad sensor. When kept in the analysis, an artifact present on a single bad sensor can spread to other sensors, for example due to spatial projection. This is why identifying bad sensors is crucial for data cleaning techniques such as the very popular Signal Space Separation (SSS) method (Taulu et al., 2004). A common practice to mitigate this issue is to visually inspect the data using an interactive viewer and mark manually, the bad sensors and bad segments in the data. Although trained experts are very likely to agree on the annotation of bad data, their judgement is subject to fluctuations and cannot be repeated. Their judgement can also be biased due to prior training with different experimental setups or equipments, not to mention the difficulty for such experts to allocate some time to review the raw data collected everyday.

Luckily, popular software tools such as Brainstorm (Tadel et al., 2011), EEGLAB (Delorme and Makeig, 2004), FieldTrip (Oostenveld et al., 2011), MNE (Gramfort et al., 2013) or SPM (Litvak et al., 2011) already allow for the rejection of bad data segments based on simple metrics such as peak-to-peak signal amplitude differences that are compared to a manually set threshold value. When the peak-to-peak amplitude in the data exceeds a certain threshold, it is considered as bad. However, while this seems quite easy to understand and simple to use from a practitioner’s standpoint, this is not always convenient. In fact, a good peak-to-peak threshold turns out to be data specific, which means that setting it requires some amount of trial and error.

The need for better automated methods for data preprocessing is clearly shared by various research teams, as the literature of the last few years can confirm. On the one hand, are pipeline-based approaches, such as Fully Automated Statistical Thresholding for EEG artifact rejection (FASTER by Nolan et al. (2010)) which detect bad sensors as well as bad trials using fixed thresholds motivated from classical Gaussian statistics. Methods such as PREP (Bigdely-Shamlo et al., 2015), on the other hand, aim to detect and clean the bad sensors only. Unfortunately, they do not offer any solution to reject bad trials. Other methods are available to solve this problem. For example, the Riemannian Potato (Barachant et al., 2013) technique can identify the bad trials as those where the covariance matrix lies outside of the “potato” of covariance matrices for good trials. By doing so, it marks trials as bad but does not identify the sensors causing the problem, hence not offering the ability to repair them. It appears that practitioners are left to choose between different methods to reject trials or repair sensors, whereas they are in fact intricately related problems and must be dealt with together. Robust regression (Diedrichsen and Shadmehr, 2005) also deals with bad trials using a weighted average which mitigates the effect of outlier trials. Trials with artifacts end up with low contributions in the average. This however does not help when analyses have to be conducted on single trials. A somewhat related approach, which is also data-driven, is Sensor Noise Suppression (SNS) (De Cheveigné and Simon, 2008). It removes the sensor-level noise by spatially projecting the data of each sensor onto the subspace spanned by the principal components of all the other sensors. This projection is repeated in leave-one-sensor-out iterations so as to eventually clean all the sensors. In most of these methods, however,

there are parameters which are somewhat dataset dependent and must therefore be manually tuned.

We therefore face the same problem in automated methods as in the case of semi-automated methods such as peak-to-peak rejection thresholds, namely the tuning of model parameters. In fact, setting the model parameters is even more challenging in some of the methods when they do not directly translate into human-interpretable physical units.

This led us to adopt a pragmatic approach in terms of algorithm design, as it focuses on the tuning of the parameters that M/EEG users presently choose manually. The goal is, not only to obtain high quality data but also to develop a method which is transparent and not too disruptive for the majority of M/EEG users. A first question we address below is: can we improve peak-to-peak based rejection methods by automating the process of trial and error? In the following section, we explain how the widely-known statistical method of cross-validation (see Figure 1 for a preview) in combination with Bayesian optimization (Snoek et al., 2012; Bergstra et al., 2011) can be employed to tackle the problem at hand. We then explain how this strategy can be extended to set thresholds separately for each sensor and mark trials as bad when a large majority of the sensors have high-amplitude artifacts. This process closely mimics how a human expert would mark a trial as bad during visual inspection.

In the rest of paper, we detail the internals of our algorithm, compare it against various state-of-the-art methods, and position it conceptually with respect to these different approaches. For this purpose, we make use of qualitative visualization techniques as well as quantitative reports. In a major validation effort, we take advantage of cleaned up evoked response fields (ERFs) provided by the Human Connectome Project (Larson-Prior et al., 2013) enabling ground truth comparison between alternative methods. This work represents one of the first efforts in reanalysis of the MEG data from the HCP dataset using a toolkit stack significantly different from the one employed by the HCP consortium. In addition to this, we validated our algorithm on the MNE sample data (Gramfort et al., 2013), the multimodal faces dataset (Wakeman and Henson, 2015), and the EEGBCI motor imagery data (Goldberger et al., 2000; Schalk et al., 2004).

A preliminary version of this work was presented in Jas et al. (2016).

2 Materials and methods

We will first describe how a cross-validation procedure can be used to set peak-to-peak rejection thresholds globally (*i.e.* same threshold for all sensors). This is what we call *autoreject (global)*.

2.1 Autoreject (global)

We denote the data matrix by $X \in \mathbb{R}^{N \times P}$ with N trials and P features. These P features are the Q sensor-level time series, each of length T concatenated along the second dimension of the data matrix, such that $P = QT$. We divide the data into K folds (along the first dimension) with training set indices $\mathcal{T}_k$ and validation set indices $\mathcal{V}_k = [1..N] \setminus \mathcal{T}_k$ for each fold k ($1 \leq k \leq K$). For simplicity of notation, we first define the peak-to-peak amplitude for the i th trial and j th sensor as the difference between the maximum and the minimum value in that time series:

$$\mathcal{A}_{ij} = \max_{(j-1)T+1 \leq t \leq jT} (X_{it}) - \min_{(j-1)T+1 \leq t \leq jT} (X_{it}) . \quad (1)$$

The good trials $\mathcal{G}_k$ where the peak-to-peak amplitude $\mathcal{A}_{ij}$ for any sensor does not exceed the candidate threshold τ are computed as

$$\mathcal{G}_k = \{i \in \mathcal{T}_k \mid \max_{1 \leq j \leq Q} \mathcal{A}_{ij} \leq \tau\}. \quad (2)$$

By comparing the peak-to-peak threshold with the maximum of the peak-to-peak amplitudes, we ensure that none of the sensors exceed the given threshold. Once we have applied the threshold on the training set, it is necessary to evaluate how the threshold performs by looking at new data. For this purpose, we consider the validation set. Concretely speaking, we propose to compare the mean $\overline{X_{\mathcal{G}_k}}$ of good trials in the training set against the median $\widetilde{X_{\mathcal{V}_k}}$ of all trials in the validation set. Using root mean squared error (RMSE) the mismatch $e_k(\tau)$ reads as:

$$e_k(\tau) = \|\overline{X_{\mathcal{G}_k}} - \widetilde{X_{\mathcal{V}_k}}\|_{\text{Fro}}. \quad (3)$$

Here, $\|\cdot\|_{\text{Fro}}$ is the Frobenius norm. The rationale for using the median in the validation set is that it is robust to outliers. Indeed, it is far less affected by high-amplitude artifacts than a classical mean. The threshold with the best data quality (lowest mismatch $e_k(\tau)$) on average across the K folds is selected as the optimal threshold. In practice τ is taken in a bounded interval $[\tau_{\min}, \tau_{\max}]$:

$$\tau_{\star} = \underset{\tau \in [\tau_{\min}, \tau_{\max}]}{\operatorname{argmin}} \frac{1}{K} \sum_{k=1}^K e_k(\tau) \quad (4)$$

Note, that $\widetilde{X_{\mathcal{V}_k}}$ does not depend on τ . Indeed, it would not be wise to restrict the validation set to good trials according to the value of τ . As τ varies, it would lead to a variable number of validation trials, which would affect the comparison of RMSE across threshold values. The idea of using the median in the context of cross-validation has been previously proposed in the statistics literature in order to deal also with outliers (Leung, 2005; De Brabanter et al., 2003).

Figure 1 shows how the average RMSE changes as the threshold varies for the MNE sample dataset (Gramfort et al., 2013, 2014). At low thresholds, our model underfits as it drops most of the trials in the data resulting in a noisy average. On the other hand, at high thresholds, the model overfits retaining all the trials in the data including the high-amplitude artifacts. Here the candidate values of τ were taken on a grid. More details on how to solve (4) will be given in Section 2.3.

2.2 Autoreject (local)

A global threshold common to all sensors, however, suffers from limitations. A common case of failure is when a single sensor is affected (locally or globally) by high-amplitude artifacts. In this case, $\max_j \mathcal{A}_{ij}$, which would be the peak-to-peak amplitude that is compared to the threshold, comes from this bad sensor. If the sensor is not repaired or removed, we might end up rejecting a large fraction of otherwise good trials, just because of a single bad sensor. This is certainly not optimal. In fact, a possibly better solution is to replace the corrupted signal in the sensor by the interpolation of the signals in the nearby sensors. A second observation is that sensors can have very different ranges of amplitudes depending on their location on the scalp. A threshold tuned for one sensor may not work as effectively for another sensor. Both of these observations are motivations for estimating rejection thresholds for each sensor separately.

Once we define sensor-wise rejection thresholds $\tau_{\star}^j$, we can define an indicator matrix $C_{ij} \in \{0, 1\}^{N \times Q}$

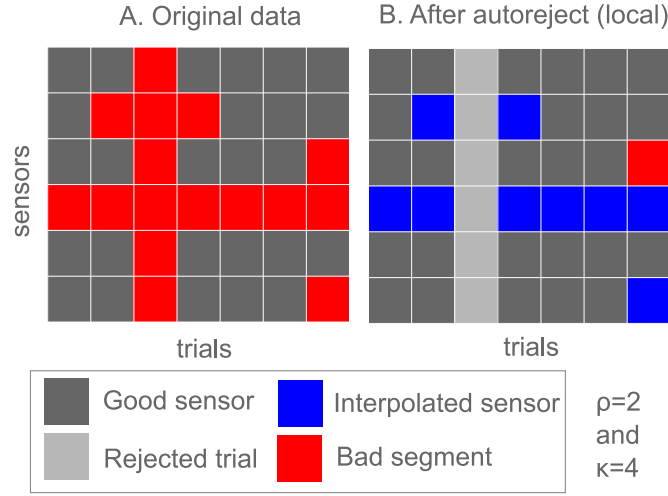


Fig. 2: A schematic diagram explaining how *autoreject (local)* works. (A) Each cell here is an element of the indicator matrix C_{ij} described in Section 2.2. Sensor-level thresholds are found and bad segments are marked for each sensor. Bad segments shown in red are where $C_{ij} = 1$ (B) Trials are rejected if the number of bad sensors is greater than κ and otherwise, the worst ρ sensors are interpolated.

which designates the bad trials at the level of individual sensors. In other words, we have:

$$C_{ij} = \begin{cases} 0, & \text{if } \mathcal{A}_{ij} \leq \tau_{\star}^j \\ 1, & \text{if } \mathcal{A}_{ij} > \tau_{\star}^j \end{cases} \quad (5)$$

The schematic in Figure 2A shows a cartoon figure for this indicator matrix C_{ij} . Now that we have identified bad sensors for each trial, one might be tempted to interpolate all the bad sensors in each trial. However, it is not as straightforward since in some trials, a majority of the sensors may be bad. These trials cannot be repaired by interpolation and must be rejected. In some other cases, the number of bad sensors may not be large enough to justify rejecting the trial. However, it might already be too much to interpolate all the sensors reliably. In these cases, a natural idea is to pick the worst few sensors and interpolate them. This suggests an algorithm as described in Figure 2B. Reject a trial only if most sensors “agree” that the trial is bad, otherwise interpolate as many sensors as possible. We will denote by κ the maximum number of bad sensors in a non-rejected trial and by ρ the maximum number of sensors that can be interpolated. Note that ρ is necessarily less than κ . The interpolation scheme for EEG uses spherical splines (Perrin et al., 1989) while for MEG it uses a Minimum Norm Estimates formulation with spherical harmonics (Hämäläinen and Ilmoniemi, 1994). The implementation is provided by MNE-Python (Gramfort et al., 2013).

The set of good trials $\mathcal{G}_k^{\kappa}$ in the training set $\mathcal{T}_k$ can therefore be written mathematically as:

$$\mathcal{G}_k^{\kappa} = \{i \in \mathcal{T}_k \mid \sum_{j=1}^Q C_{ij} < \kappa\} . \quad (6)$$

In the remaining trials, if $\rho < \kappa$, one needs to define what are the worse ρ sensors that shall be interpolated. To do this we propose to rank the sensors for “badness” according to a score. A natural strategy to set the score is to use the peak-to-peak amplitude itself:

$$s_{ij} = \begin{cases} \mathcal{A}_{ij} & \text{if } C_{ij} = 1 \\ -\infty & \text{if } C_{ij} = 0 \end{cases} \quad (7)$$

Higher the score s_{ij} , the worse is the sensor. The $-\infty$ score is for ignoring the good sensors in the subsequent step. The following strategy is used for interpolation. If the number of bad sensors $\sum_{j'=1}^Q C_{ij'}$ is less than ρ we will interpolate all of them. Otherwise, we will interpolate the ρ sensors with the highest scores. In other words, we interpolate at most $\min(\rho, \sum_{j'=1}^Q C_{ij'})$ sensors.

Denoting by $X_{\mathcal{G}_k}^\rho$ the data in the training set after rejection and cleaning by interpolation, the RMSE averaged over K folds for the parameter pair (ρ, κ) therefore becomes:

$$\bar{e}(\rho, \kappa) = \frac{1}{K} \sum_{k=1}^K \|\overline{X_{\mathcal{G}_k}^\rho} - \widetilde{X_{\mathcal{V}_k}}\|_{\text{Fro}} \quad (8)$$

where $\|\cdot\|_{\text{Fro}}$ is the Frobenius norm. Finally, the best parameters ρ_* and κ_* are estimated using grid search (Hsu et al., 2003).

2.2.1 Data augmentation

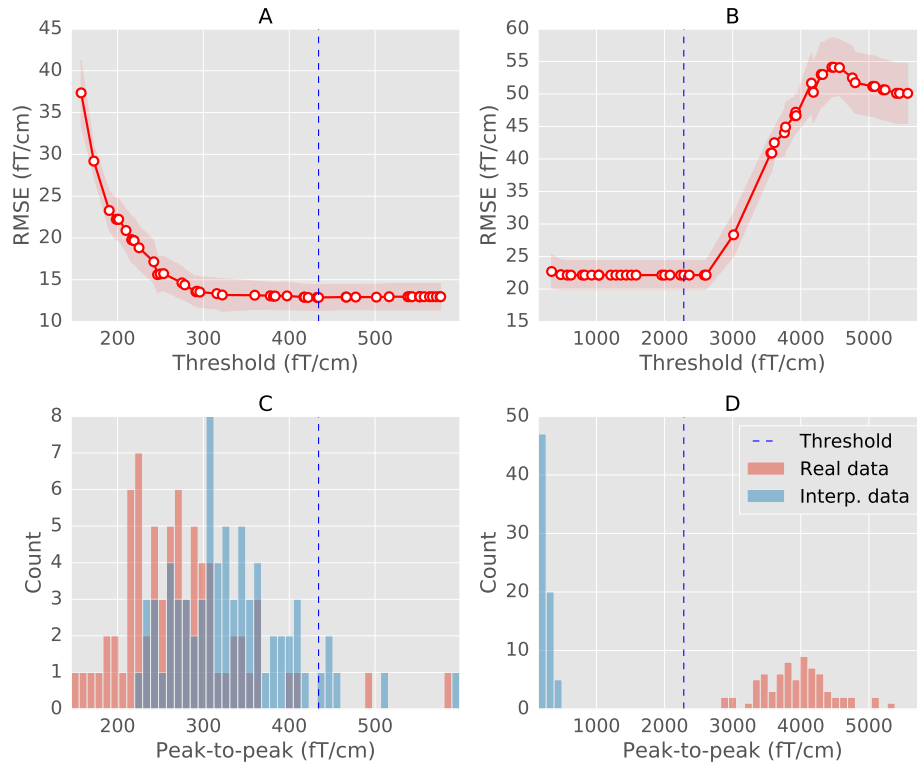


Fig. 3: (A) and (B) The cross-validation curve obtained with sequential Bayesian optimization (see Section 2.3 for an explanation) for a regular (MEG 2523) and a globally bad sensor (MEG 2443) from the MNE sample dataset. The mean RMSE is shown in red circles with error bounds in red shades. The red shaded region shows the lower and upper bounds between which the optimization is carried out. Vertical dashed line marks the estimated threshold. (C) and (D) Histogram of peak-to-peak amplitudes of trials in the sensor. The histograms are computed separately for the real data (red) and the data interpolated from other sensors (blue). The estimated threshold correctly marks all the trials as bad for the globally bad sensor.

In practice, cross-validation does not work for a globally bad sensor since all the trials are corrupted.

In this scenario, the optimal threshold for this bad sensor should be lower than the lowest peak-to-peak amplitude so that all the trials for that sensor are marked as bad. However, even the median of the validation set has been corrupted. The algorithm therefore attempts to keep as many trials as necessary for the average to be close to the corrupted median. Thus, the estimated threshold ends up being higher than what would have been optimal. Recall from Figure 1 that this is the classic case of an overfitting model. A common strategy in machine learning to reduce overfitting is data augmentation (Krizhevsky et al., 2012). It basically boils down to using the properties of the data, such as the physics of the system, to generate more plausible data.

To implement data augmentation in our model, we interpolate each sensor from all the other $Q - 1$ sensors and by doing so, we double the number of trials in the data. In the augmented data, half of the trials contain sensor data which are the output of a leave-one-sensor-out cross-validation. The augmented data matrix is $X^{\text{aug}} \in \mathbb{R}^{2N \times P}$. With the augmented data, the median is now closer to the uncorrupted median of the data in that sensor. During cross-validation the folds were stratified so that the number of interpolated trials and original trials in each fold were roughly equal.

2.3 Candidate thresholds using Bayesian optimization

Now that we have formalized the problem and our approach, we must estimate the threshold τ_* which minimizes the error defined in Equation (3). A naïve strategy is to define a set of equally spaced points over a range of thresholds $[\tau_{\min}, \tau_{\max}]$. The estimated threshold would be the one which obtains the lowest error among these candidate threshold. This is the approach taken in Figure 1. The range of thresholds is easy to set as it can be determined from the minimum and maximum peak-to-peak amplitude for the sensor in the augmented data matrix X^{aug} . However, it is not obvious how to set the spacing between the candidate thresholds, and experiments showed that varying this spacing could impact the results. If the candidate thresholds are far apart, one might end up missing the optimal threshold. On the other hand, if the thresholds are very dense, it is computationally more demanding.

This motivated us to use Bayesian optimization (Snoek et al., 2012; Bergstra et al., 2011) to estimate the optimal thresholds. It is a sequential approach which decides the next candidate threshold to try based on all the observed thresholds so far. It is based on maximizing an acquisition function given an objective function of samples seen so far (data likelihood) and the prior (typically a Gaussian Process (GP) (Rasmussen and Williams, 2006)). The objective function in our case is the mean cross-validation error as defined in Equations (3). To obtain the next iterate, an acquisition function is maximized over the posterior distribution. Popular choices of the acquisition function include “probability of improvement”, “expected improvement” and “confidence bounds of the GP” (Snoek et al., 2012). We pick “expected improvement” as it balances exploration (searching unknown regions) and exploitation (maximizing the improvement) strategies without the need of a tuning parameter. For our analysis, we use the scikit-optimize¹ implementation of Bayesian optimization, which internally uses the Gaussian process module from scikit-learn (Pedregosa et al., 2011).

Figure 3A and 3B show the cross-validation curve for a regular sensor and a globally bad sensor in the MNE sample dataset (Gramfort et al., 2014, 2013). The RMSE is evaluated on thresholds as determined by the Bayesian optimization rather than a uniform grid. These plots also illustrate the arguments presented in Section 2.2.1 with respect to data augmentation. The histograms in Figure 3C for the interpolated data and

¹ <https://scikit-optimize.github.io>

Tab. 1: Overview of rejection strategies evaluated

method	statistical scope	parameter defaults
FASTER	univariate	threshold on zscore = 3
SNS	multivariate	number of neighbors = 8
RANSAC	multivariate outlier detection	#resamples = 50, fraction of channels = 0.25, threshold on correlation = 0.75, unbroken time = 0.4
autoreject	univariate with cross-validation	sensor-level thresholds, ρ and κ ; learned from data

the real data are overlapping for the regular sensor. Thus, the estimated threshold for that sensor marks a trial as outlier if its peak-to-peak values is much higher than the rest of the trials. However, in the case of a globally bad sensor, the histogram (Figure 3D) is bimodal – one mode for the interpolated data and one mode for the real data. Now, the estimated threshold is no longer marking outliers in the traditional sense. Instead, all the trials belonging to that sensor must be marked as bad.

3 Experimental Validation Protocol

To experimentally validate *autoreject*, our general strategy is to first visually evaluate the results and thereafter quantify the performance. We describe below the evaluation metric used, the methods we compare against, and finally the datasets analyzed. All general data processing was done using the open source software MNE-Python (Gramfort et al., 2013).

3.1 Evaluation metric

The evoked response from the data cleaned using our algorithm or a competing benchmark is denoted by $\bar{X}(\text{method})$. This is compared to the ground truth evoked response $\bar{X}(\text{clean})$ (See Section 3.2.1 to see how these are obtained for different datasets) using:

$$\|\bar{X}(\text{method}) - \bar{X}(\text{clean})\|_{\infty} \quad (9)$$

where $\|\cdot\|_{\infty}$ is the infinity norm. The reason for using infinity norm is that it is sensitive to the maximum amplitude in the difference signal as opposed to the Frobenius norm which averages the squared difference. The $\|\cdot\|_{\infty}$ is a particularly sensitive metric to quantify artifacts which are also visually striking such as those localized on one sensor or at a given time instant.

3.2 Competing methods

Here, we list the methods that will be quantitatively compared to *autoreject* using the evaluation metric in Equation 9. These methods are also summarized for the reader’s convenience in Table 1.

- *No rejection*: It is a simple sanity check to make sure that the data quality upon applying the *autoreject* (*local*) algorithm does indeed improve. This is the data before the algorithm is applied.
- *Sensor Noise Suppression (SNS)*: The SNS (De Cheveigné and Simon, 2008) algorithm, as described in the Introduction (Section 1), projects the data of each sensor on to the subspace spanned by the principle components of all the other sensors. What it does is regressing out the sensor noise that cannot be explained by other sensors. It works on the principle that brain sources project on to

Tab. 2: Overview of datasets analyzed

Algorithm	Dataset	Acquisition device	Sensors used	#subjects
autoreject (global)	MNE sample data	Neuromag VectorView	60 EEG electrodes	1
	EEGBCI	BCI2000 cap	64 EEG electrodes	105
autoreject (local)	MNE sample data	Neuromag VectorView	60 EEG electrodes	1
	EEG faces	Neuromag VectorView	60 EEG electrodes	19
	HCP working memory	4D Magnes 3600 WH	248 magnetometers	83

multiple sensors but the noise is uncorrelated across sensors. In practice, not all the sensors are used for projection, but only a certain number of neighboring sensors (determined by the correlation in the data between the sensors).

- *Fully Automated Statistical Thresholding for EEG artifact Rejection (FASTER)*: It finds the outlier sensor using five different criteria: the variance, correlation, hurst, kurtosis and line noise. When the z-score of any of these criteria exceeds 3, the sensor is marked as bad according to that criteria. Note that even though FASTER is typically used as an integrated pipeline, here we use the bad sensor detection step, as this is what appears to dominate the bad signals in the case of the HCP data (Section 3.2.1). We take a union of the sensors marked as bad by the different criteria and interpolate the data for those sensors.
- *Random Sample Consensus (RANSAC)*: We use the RANSAC implemented as part of the PREP pipeline (Bigdely-Shamlo et al., 2015). In fact, RANSAC (Fischler and Bolles, 1981) is a well-known approach used to fit statistical models in the presence of outliers in the data. In this approach, adopted for the use case of artifact detection in EEG, a subset of sensors (inliers) are sampled randomly (25% of the total sensors) and the data in all sensors are interpolated from these inliers sensors. This is repeated multiple times (50 in the PREP implementation) so as to yield a set of 50 time series for each sensor. The correlation between the median, computed instant by instant, of these 50 time series and the real data is computed. If this correlation is less than a threshold (0.75 in the PREP implementation), then the sensor is considered an outlier and therefore marked as bad. It is perhaps worth noting that unlike in the classical RANSAC algorithm, the inlier model is not learned from the data but instead determined from the physical interpolation. A sensor which is bad for more than 40% of the trials (the unbroken time) is marked as globally bad and interpolated. Even though the method was first proposed on EEG data only, we extended it for MEG data by replacing spline interpolation with field interpolation using spherical harmonics as implemented in MNE (Gramfort et al., 2013; Hämäläinen and Ilmoniemi, 1994). Note that this is the same interpolation method that is used by *autoreject (local)*.

3.2.1 Datasets

We validated our methods on four open datasets with data from over 200 subjects. This allowed us to evaluate experimentally strengths and potential limitations of different rejection methods. The datasets contained either EEG or MEG data. To obtain solid experimental conclusions, diverse experimental paradigms were considered with data from working memory, perceptual and motor tasks.

We detail below how we defined $\bar{X}(clean)$, the cleaned ground-truth data for two of our datasets – HCP

MEG and EEG faces data. This is perhaps one of the most challenging aspects of this work because the performance is evaluated on real data and not on simulations. An overview of all the datasets used in this study is provided in Table 2.

MNE sample data The MNE sample data (Gramfort et al., 2013) is a multimodal open dataset consisting of MEG and EEG data. It has been integrated as the default testing dataset into the development of the MNE software (Gramfort et al., 2013). The simultaneous M/EEG data were recorded at the Martinos Center of Massachusetts General Hospital. The MEG data with a Neuromag VectorView system, and an MEG-compatible cap comprising 60 electrodes was used for the EEG recordings. Data were sampled at 600 Hz. In the experiment, auditory stimuli (delivered monaurally to the left or right ear) and visual stimuli (shown in the left or right visual hemifield) were presented in a random sequence with a stimulus onset asynchrony of 750 ms. The data was low-pass filtered at 40 Hz. The trials were 700 ms long including a 200 ms baseline period which was used for baseline correction.

EEGBCI dataset This is a 109-subject dataset (of which we analyzed 105 subjects which can be easily downloaded and analyzed using MNE-Python (Gramfort et al., 2013)) containing EEG data recording with a 64-sensor BCI2000 EEG cap (Schalk et al., 2004). Subjects were asked to perform different motor/imagery tasks while their EEG activity was recorded. In the related BCI protocol, each subject performed 14 runs, amounting to a total of 180 trials for hands and feet movements (90 trials each). The data was band-pass filtered between 1 and 40 Hz, and 700 ms long trials were constructed including a 200 ms pre-stimulus baseline period.

EEG faces data (OpenfMRI ds000117) The OpenfMRI ds000117 dataset (Wakeman and Henson, 2015) contains multimodal task-related neuroimaging data over multiple runs for EEG, MEG and fMRI. For our analysis, we restrict ourselves to EEG data. The EEG data was recorded using a 70 channel Easycap EEG with electrode layout conforming to the 10-10% system. Subjects were presented with images of famous faces, unfamiliar faces and scrambled faces as stimuli. For each subject, on average, about 293 trials were available for famous and unfamiliar faces. The authors kindly provided us with run-wise bad sensor annotations which allowed us to conduct benchmarking against human judgement. To generate the ground truth evoked response $\bar{X}(clean)$, we randomly select 80 percent of the total number of trials in which famous and unfamiliar faces were displayed. In these trials, we interpolated the bad sensors run-wise. Then, we removed physiological artifacts (heart beat and eye blinks) using Independent Component Analysis (ICA) (Vigário et al., 2000). Following the ICA pipelines recommended by the MNE-Python software, the bad ICA components were marked automatically using cross-trial phase statistics (Dammers et al., 2008) for ECG (threshold=0.8) and adaptive z-scoring (threshold=3) for EOG components. The evoked response from the cleaned data $\bar{X}(method)$ is computed from the remaining 20 percent trials cleaned using either *autoreject* (*local*) or *RANSAC* (see Section 4.3 for a description of this method). Computing the ground-truth evoked potential from a large proportion of trials minimized the effect of outliers in the average. However, it is noteworthy that this choice of assigning fewer trials to the estimation with rejection algorithms acts in a conservative sense: each unnoticed bad trial may affect the ensuing evoked potentials more severely.

Human Connectome Project (HCP) MEG data The HCP dataset is a multimodal reference dataset realized by the efforts of multiple international laboratories around the world. It currently provides access to

both task-free and task-related data for more than 900 human subjects with functional MRI data, 95 of which have presently also MEG (Larson-Prior et al., 2013). An interesting aspect of the initiative is that the data provided is not only in unprocessed BTi format, but also processed using diverse processing pipelines. These include annotations of bad sensors and corrupted time segments for the MEG data derived from automated pipelines and supplemented by human inspection. The automated pipelines are based on correlation between neighboring sensors, z-score metrics, ratio of variance to neighbors, and ICA decomposition. Most significant for our purposes, the clean average response $\bar{X}(clean)$ is directly available. It allows us to objectively evaluate the proposed algorithm against state-of-the-art methods by reprocessing the raw data and comparing the outcome with the official pipeline output.

The HCP MEG dataset provides access to MEG recordings from diverse tasks, *i.e.*, a motor paradigm, passive listening and working memory. Here, we focused on the working memory task for which data is available for 83 subjects out of 95. A considerable proportion of subjects were genetically related, but we can ignore this information as the purpose of our algorithm is artifact removal rather than analyzing brain responses. For each subject two runs are available. Two classes of stimuli were employed, faces and tools. Here, we focused on the MEG data in response to stimulus onsets for the “faces” condition.

The MEG data were recorded with a wholehead MAGNES 3600 (4D Neuroimaging, San Diego, CA) in a magnetically shielded room at Saint Louis University. The system comprises 248 magnetometers and 23 reference sensors to capture environmental signals. Time windows precisely matched values used by the HCP “eravg” pipeline with onsets and offsets at -1.5 s and 2.5 s before and after the stimulus event, respectively. As in the HCP pipeline, signals were down-sampled to 508.63 Hz and band-pass filtered between 0.5–60 Hz. As it is commonly done with BTi systems, reference sensors at the periphery of the head were used to subtract away environmental noise. Given the linearity of Maxwell equations in the quasi-static regime, a linear regression model was employed. More precisely, signals from reference sensors are used as regressors in order to predict the MEG data of interest. The ensuing signal explained by the reference sensors in this model was then removed. The HCP preprocessing pipeline contains two additional steps: ICA was used to remove components not related to brain activity (including eye blinks and heart beats) and then bad trials and bad segments were removed with a combination of automated methods as well as annotations by a human observer. To have a fair comparison and focus on the latter step, the ICA matrices provided by the HCP consortium were applied to the data. We interpolated the missing sensors in $\bar{X}(clean)$ so that it has the same dimensions as the data from $\bar{X}(method)$. All the algorithms were executed separately on each run and the evoked response of the two runs was averaged to get $\bar{X}(method)$.

To enable easy access of the files along with compatibility in MNE-Python, we make use of the open source MNE-HCP package². For further details on the HCP pipelines, the interested reader can consult the related paper by Larson-Prior et al. (2013) and the HCP S900 reference manual for the MEG3 release.

4 Results

We conducted qualitative and quantitative performance evaluations of *autoreject* using four different datasets comparing it to a baseline condition without rejection as well as three different alternative artifact rejection procedures.

² <http://mne-tools.github.io/mne-hcp/>

4.1 Peak-to-peak thresholds

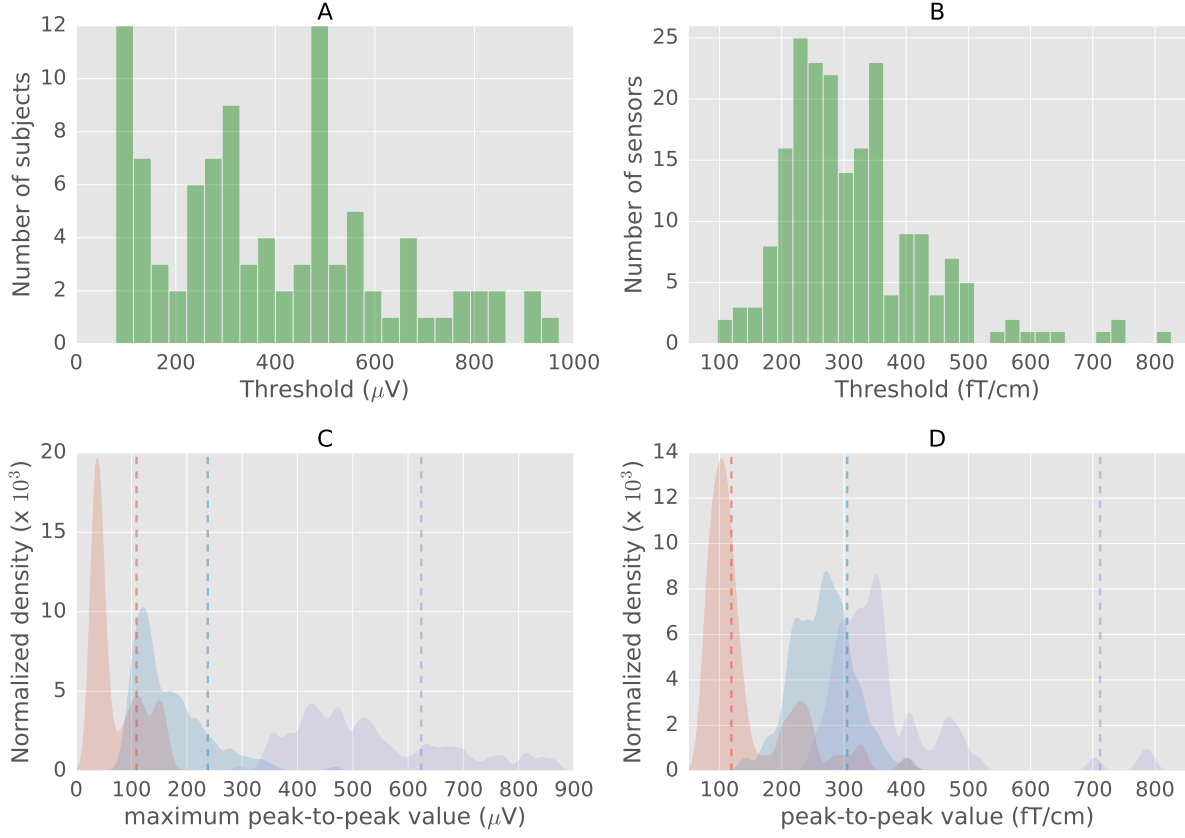


Fig. 4: A. Histogram of thresholds for subjects in the EEGBCI dataset with *autoreject (global)* B. Histogram of sensor-specific thresholds in gradiometers for the MNE sample dataset (Section 4). C. Normalized kernel density plots of maximum peak-to-peak value across sensors for three subjects in the EEGBCI data. Vertical dashed lines indicate estimated thresholds. Density plots and thresholds corresponding to the same subject are the same color. D. Normalized Kernel Density plots of peak-to-peak values for three MEG sensors in the MNE sample dataset. The threshold indeed has to be different depending on the data (subject and sensor).

First, let us convince ourselves that the peak-to-peak thresholds indeed need to be learned. In Figure 4A, we show a histogram of the thresholds learned on subjects in the EEGBCI dataset using *autoreject (global)*. This figure shows that thresholds vary a lot across subjects. One could argue that this is due to variance in the estimation process. To rule out such a possibility, we plotted the distribution of maximum peak-to-peak thresholds as kernel density plots in Figure 4C for three different subjects. We can see that these distributions are indeed subject dependent, which is why a different threshold must be learned for each subject. In fact, if we were to use a constant threshold of $150\mu V$, in 17% of the subjects, all the trials would be dropped in one of the two conditions. Of course, from Figure 4A, we can now observe that $150\mu V$ is not really a good threshold to choose for many subjects.

We show here the maximum peak-to-peak amplitude per sensor because this is what decides if a trials should be dropped or not in the case of *autoreject (global)*. Note that, if instead, we examined the distribution of peak-to-peak amplitudes across all sensors and trials, we would see a quasi-normal distribution. When all

the sensors are taken together, a “smoothing” effect is observed in the distribution. This is a consequence of the central limit theorem. This also explains why we cannot learn a global threshold using all the peak-to-peak amplitudes across trials and sensors.

With the *autoreject* (*local*) approach, a threshold is estimated for each sensor separately. The histogram of thresholds for the MNE sample dataset is plotted in Figure 4B. It shows that the threshold varies even across homogeneous MEG sensors. Figure 4D shows the distribution of peak-to-peak thresholds for three different MEG sensors. This graph confirms actual sensor-level differences in amplitude distributions, which was also previously reported in the literature (Junghöfer et al., 2000). With this work, we go one step further by learning automatically the thresholds in a data-driven way rather than asking users to mark them interactively.

4.2 Visual quality check

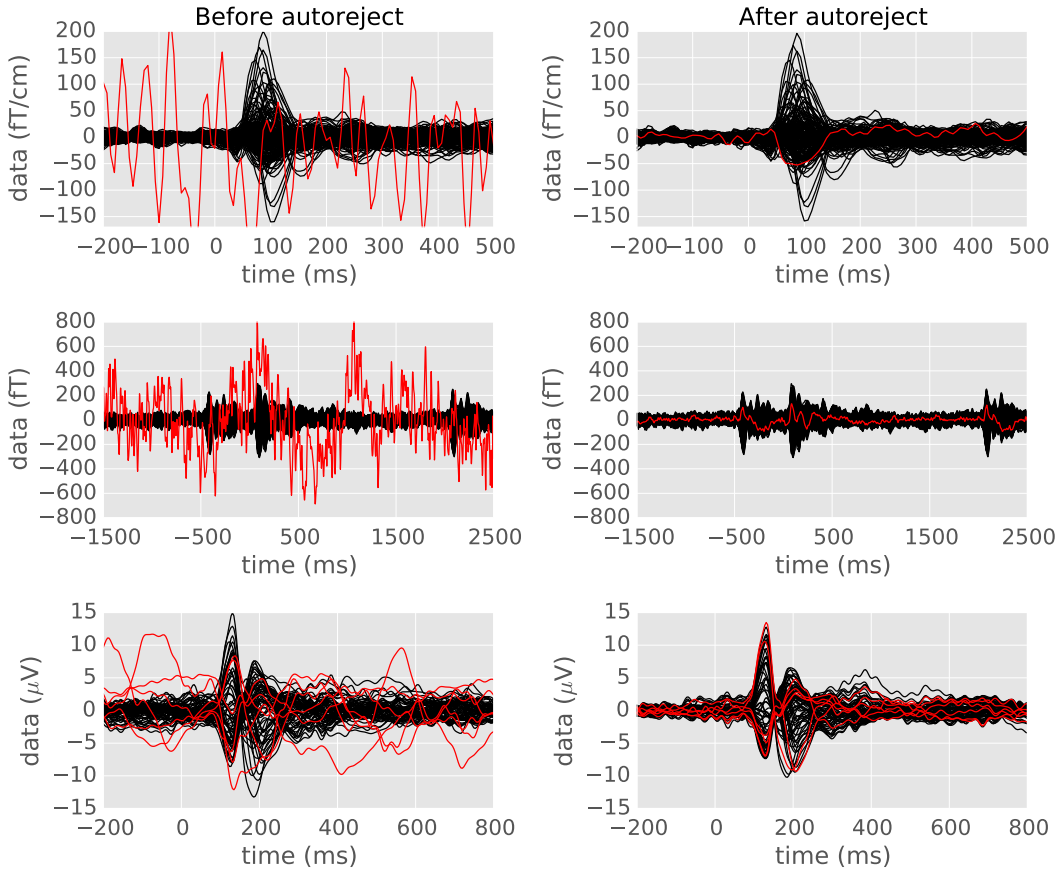


Fig. 5: The evoked response (average of data across trials) on three different datasets before and after applying *autoreject* — the MNE sample data, the HCP data and the EEG faces data. Each sensor is a line on the plots. On the left, manually annotated bad sensors are shown in red. The algorithm finds the bad sensors automatically and repairs them for the relevant trials. Note that it can even fix multiple sensors at a time and works for different modalities of data acquisition.

The average response plotted in a single graph, better known as “butterfly plots”, constitutes a natural way to visually assess the performance of the algorithm for three different datasets – MNE sample data, HCP

MEG data, and EEG faces data. In Figure 5, the subplots in the left column show the evoked response with the bad sensors marked in red. Right subplots, show data after applying the *autoreject (local)* algorithm, with the repaired bad sensors in red. The algorithm works for different acquisition modalities – MEG and EEG, and even when multiple sensors are bad. A careful look at the results, show that *autoreject (local)* does not completely remove eyeblinks in the data as some of the blinks are time-locked to the evoked response. We will later discuss (Section 5) the possible solutions of applying ICA-based artifact correction in combination with *autoreject (local)*.

4.3 Quantification of performance and comparison with state-of-the-art

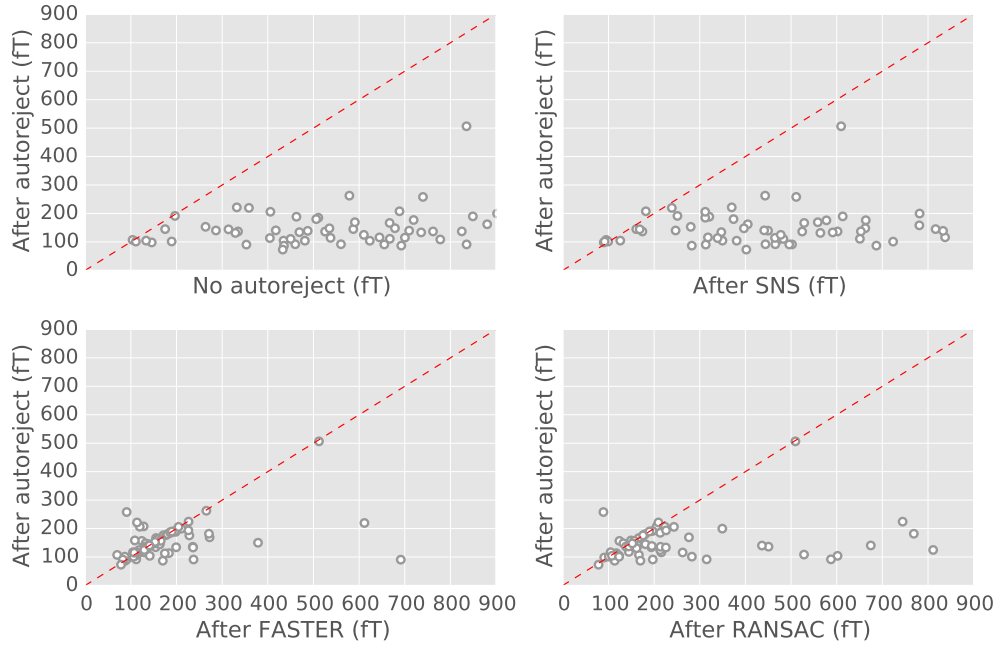


Fig. 6: Scatter plots for the results with the HCP data. For each method, the $\|\cdot\|_\infty$ norm of the difference between the HCP ground truth and the method is taken. Each circle is a subject. (A) *autoreject (local)* against no rejection, (B) *autoreject (local)* against Sensor Noise Suppression (SNS) (SNS), (C) *autoreject* against FASTER, (D) *autoreject (local)* against RANSAC. Data points below the dotted red line indicate subjects for which *autoreject (local)* outperforms the alternative method.

We now compare these algorithms to *autoreject (local)* using the data quality metric defined in Equation (9). We are interested not only in how the algorithms perform on average but at the level of individual subjects. To detail single subject performance, we present the data quality as scatter plots where each axis corresponds to the performance of a method. Figure 6, contains results on the HCP MEG data. We can observe from the top-left subplot of the figure that *autoreject (local)* does indeed improve the data quality in comparison to the *no rejection* approach. In Figure 6B, *autoreject (local)* is compared against SNS. The SNS algorithm focuses on removing noise isolated on single sensors. Its results can be affected by the presence of multiple bad sensors and globally bad trials. This explains why *autoreject (local)* outperforms SNS in this setting. In Figure 6C, we compare against FASTER. Even though *autoreject (local)* is slightly worse than FASTER for a few subjects, FASTER is clearly much worse than *autoreject (local)* for at least 3 subjects, and

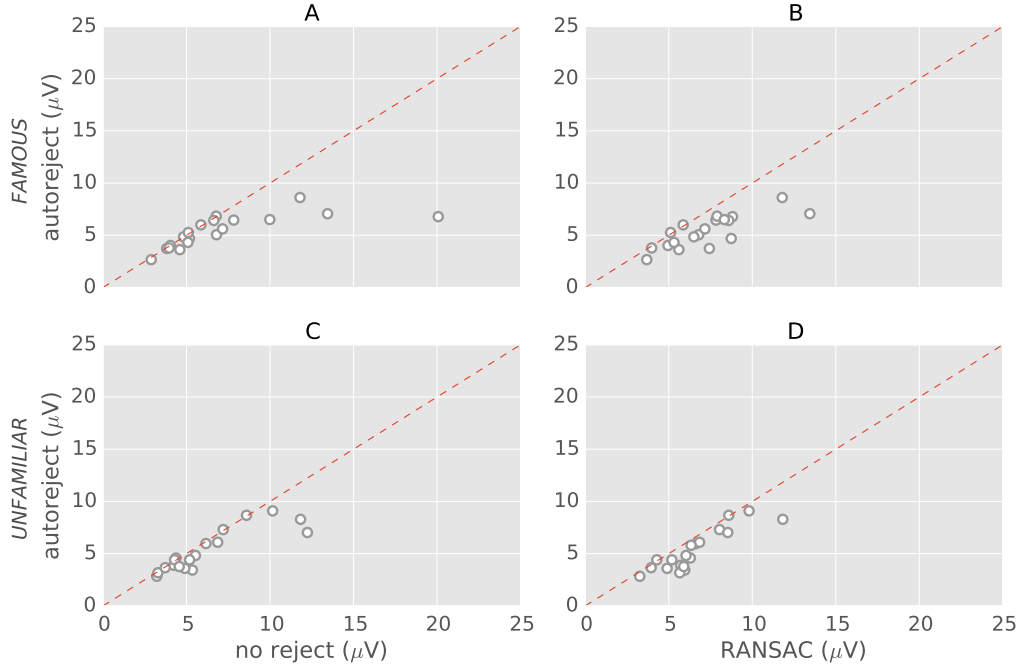


Fig. 7: Scatter plots for the results with the 19 subjects from Faces dataset. The first row (A) and (B) is for the condition “famous” and the second row (C) and (D) is for the condition “unfamiliar” faces. For each method, the $\|\cdot\|_\infty$ norm of the difference between the ground truth and the estimates is computed. Each circle is a subject. Data points below the dotted red line indicate subjects for which *autoreject (local)* outperforms the alternative method.

autoreject (local) yields therefore less errors on average. Finally, Figure 6D shows comparison to RANSAC. In the PREP implementation, this algorithm is not fully data-driven in the classic sense of RANSAC. This is due to the fact that the inlier model is not learned but rather derived from the physics of the interpolation. It is therefore an algorithm which is conceptually close to *autoreject*. However, a critical difference is that the parameters of this method still need to be tuned. This can be a problem as these parameters can be suboptimal on some datasets. Some experiments showed that it is for example the case for the EEG faces data, where it is possible to obtain better results by manually tuning the RANSAC parameters, rather than using the values proposed by the original authors.

Figure 7 presents scatter plots for the EEG faces data. Here, we restrict our comparison to RANSAC as it is conceptually the closest to *autoreject*. On this data, we apply the algorithms on both the conditions – famous and unfamiliar faces. It should be noted that the ground truth for this data was generated automatically with no additional annotations from human experts. However, a sanity check was performed on the ground truth by visual inspection. Here too, *autoreject* offers good results across all subjects, and even for the subjects for which RANSAC underperforms.

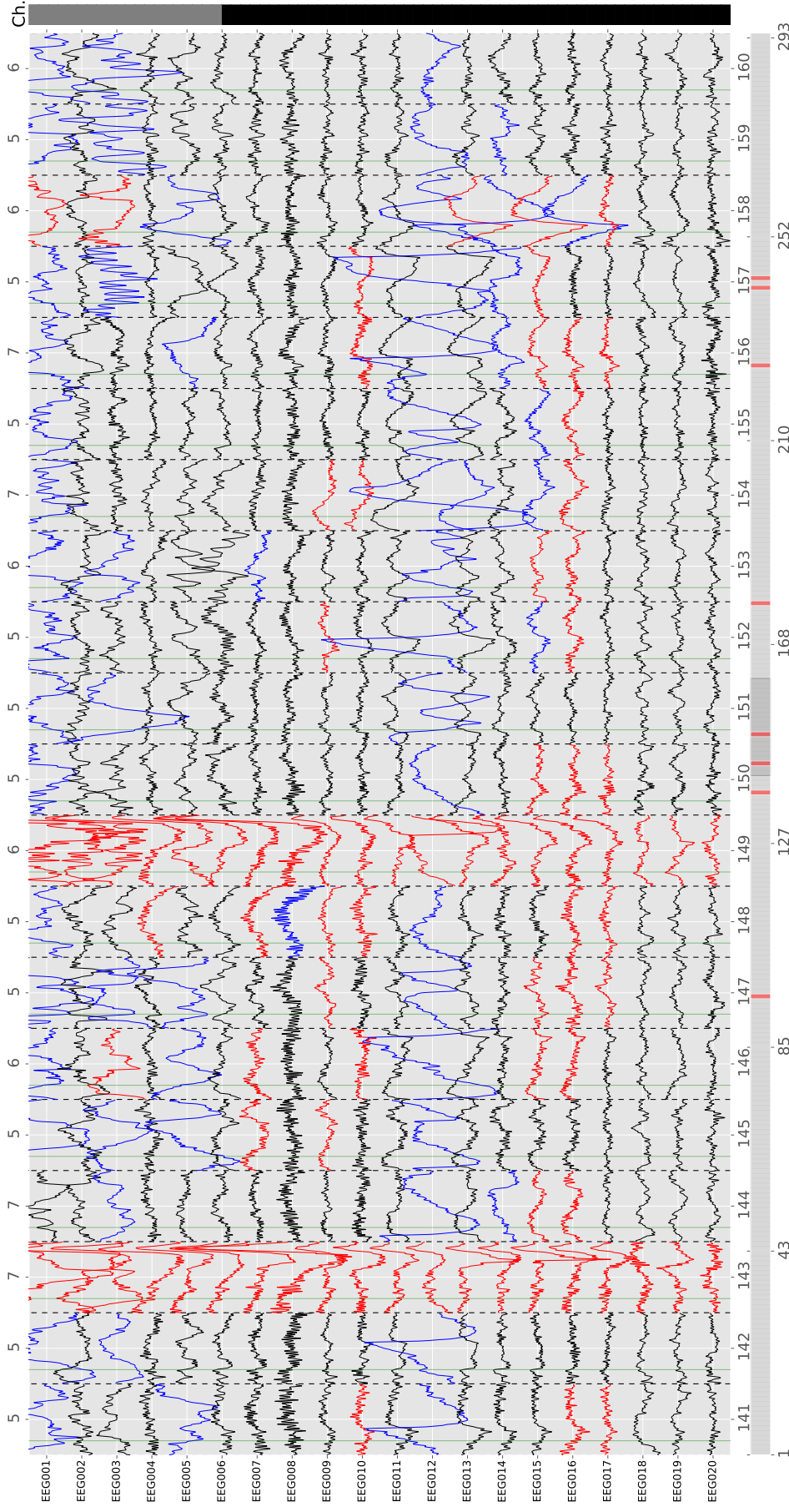


Fig. 8: An example diagnostic plot from an interactive viewer with *autoreject (local)*. The data plotted here is subject 16 for the condition ‘famous’ in the EEG faces data. Each row is a different sensor. The trials are concatenated along the x axis with dotted vertical lines separating consecutive trials. Each trial is numbered at the bottom and its corresponding trigger code is at the top. The horizontal scroll bar at the bottom allows browsing trials and the vertical scroll bar on the right is for browsing sensors. A trial which is marked as bad is shown in red on the horizontal scroll bar and the corresponding column for the trial is also red. A data segment in a good trial is either i) Good (in black) ii) Bad and interpolated (blue), or iii) Bad but not interpolated (in red). Note that the worst sensors in a trial are typically interpolated.

5 Discussion

In this study, we have presented a novel artifact rejection algorithm called *autoreject* and assessed its performance on multiple datasets showing comparisons with other methods from the state-of-the-art.

We have shown that learning peak-to-peak rejection thresholds subject-wise is justified as the distribution of this statistic indeed varies considerably across subjects. We have shown qualitatively that *autoreject* yielded clean physiological event related field (ERF) and event related potentials (ERP) by correcting or rejecting contaminated data segments. Finally, we have shown quantitatively that *autoreject* yields results closer to the ground truth for more subjects than the algorithms presented in Section 3.2. We now further discuss the conceptual similarities and differences of our approach to the alternative methods. We also discuss the interaction between *autoreject* and some other steps in the M/EEG analysis pipelines.

5.1 Autoreject vs. competing methods

We believe the key advantage of *autoreject (local)* over the other methods consists in combining data-driven parameter tuning with deterministic, physics-driven data interpolation, explicitly taking into account the well-understood Maxwell’s equations. To recapitulate, the sensor-level thresholds mark outlier segments across trials at the level of individual sensors, following a data augmentation step which exploits the full array of sensors. They are learned from the data using cross-validation. At repair time, bad segments are replaced with interpolated data from the good sensors. Of course, this is problematic if the sensor locations are not readily available. Fortunately, it turns out that the sensor positions from standard montages are often good enough for reliable interpolation.

In contrast to *autoreject (local)*, SNS is a purely statistical method that does not take into account the physics of sensor locations for repairing the data. In SNS, the sensors are considered in a leave-one-sensor-out protocol. For each sensor, a “clean” subspace is defined from the principle components of the remaining sensors. The data from this sensor is then projected on to the “clean” subspace. As we have seen in Section 4 (Figure 6), this does not work satisfactorily, presumably because the SNS method makes strong assumptions regarding the orthogonality of the noise and “clean” subspace. The ensuing projection may not improve, and even deteriorate the signal in some cases. The consequence of this is what we observe empirically in Figure 6. Applying SNS will also be problematic when multiple sensors are corrupted simultaneously. However, this is less of a problem in the HCP MEG data that we analyzed.

On the other hand, the FASTER method derives its rejection decisions from multiple signal characteristics. It uses criteria such as between-sensor correlation, variance and power spectrum, by considering their univariate Gaussian statistics with thresholds fixed to a z-score of 3. This default threshold appears to be satisfying as they work on a vast majority of subjects. However, the fact that it does not work as well on certain subjects can limit its adoption for large scale studies. Here, the adaptive nature of threshold detection performed by *autoreject* seems to be a clear advantage.

The RANSAC algorithm also performs adaptive outlier detection, but across sensors rather than trials. While, *autoreject (local)* operates on segmented data such as trials time-locked to the stimuli, RANSAC was designed for continuous data without any segmentation. In fact, one could readily obtain bad sensor per trial (as illustrated in Figure 2) even with RANSAC. However, the authors of the paper did not validate their method on continuous data, and hence, such a modification would require additional work. Although in the case of MEG data, this is not very crucial, this can in fact be critical for EEG data analysis. Remember,

that in EEG, one often has to deal with locally bad sensors. And in this context, it is noteworthy that none of the other methods we have discussed so far provides an explicit treatment for single trial analysis in the presence of locally bad sensors. Our comparison to the RANSAC algorithm seems to suggest that the RANSAC algorithm is indeed sensitive to the parameter settings. Even though the default settings appear to work reasonably well for the EEG data (Figure 7), they are not so optimal for the HCP MEG data (Figure 6).

Regardless of the method that researchers choose to adopt, diagnostic plots and automated reports (Engemann et al., 2015) are an essential element of any M/EEG analysis. In this regard, transparency of the method in question is important. In the case of our *autoreject (local)* implementation, we offer the possibility for the user to inspect the bad segments marked by the automated algorithm and correct it if necessary. An example of such a plot is shown in Figure 8. Automating the detection of bad sensors and trials has the benefit of avoiding any unintentional biases that might be introduced if the experimenter were to mark the segments manually. In this sense, diagnostic visualization should supplement the analysis by insuring accountability in the case of unexpected results.

5.2 Autoreject in the context of ICA and SSP

It is now important to place these results in the broader context of electrophysiologic data analysis. Regarding the correction of specific artifacts such as electrooculogram (EOG) artifacts, *autoreject (local)* does indeed remove or interpolate some of the trials affected by eye blinks. This is because most eye blinks are not time-locked to the trial onsets and therefore get detected in the cross-validation procedure. However, the weaker eye blinks, particularly those smaller in magnitude than the evoked response, are not always removed. Also, the idea of rejection is to remove extreme values which are supposed to be rare events. This is why our empirical observation suggests that *autoreject (local)* is not enough in the presence of too frequent eye blinks, but also not enough to fully get rid of the smallest EOG artifacts.

This is where ICA (Vigário, 1997) and Signal Space Projection (SSP) (Uusitalo and Ilmoniemi, 1997) can naturally supplement *autoreject*. These methods are usually employed to extract and subsequently project out signal subspaces governed by physiological artifacts such as muscular, cardiac and ocular artifacts. However the estimation of these subspaces can be easily corrupted by other even more dramatic environmental or device-related artifacts. This is commonly prevented by band-pass filtering the signals and excluding high-amplitude artifacts during the estimation of the subspaces. Both ICA and SSP are highly sensitive to observations with high variance. Even though they involve estimating spatial filters that do not incorporate any notion of time, artifacts very localized in time will very likely have a considerable impact on the estimation procedure. This leads us to recommend removing globally bad sensors and employing appropriate rejection thresholds to exclude time segments with strong artifacts. If the context of data processing supports estimation of ICA and SSP on segmented data, we would recommend to perform it after applying *autoreject*, benefiting from its automated bad sensor and bad trial handling.

5.3 Source localization with artifact rejection

Another common M/EEG processing step that requires clean data is the estimation of cortical or cerebral sources from M/EEG data. It is the so-called M/EEG inverse problem. As *autoreject* ends up automating an existing processing step prior to source localization, it naturally integrates with such source-level analyses. For example, evoked responses obtained using *autoreject (local)* can be readily analyzed with various source

localization methods such as beamformer methods, or cortically-constrained Minimum Norm Estimates with ℓ_2 penalty (Uutela et al., 1999), and even with its noise-normalized variants such as dSPM and sLORETA. Note also that, in contrast to denoising techniques such as SSP (Uusitalo and Ilmoniemi, 1997) or SSS (Taulu et al., 2004), *autoreject* does not systematically reduce the rank of the data, which can be problematic, in particular for beamforming techniques.

The evoked response obtained using *autoreject (local)* can be readily applied to noise-normalized source localization methods such as dSPM (Dale et al., 2000) and sLORETA (Pascual-Marqui et al., 2002). Indeed, such methods rely on the estimation of the between-sensor noise covariance during baseline periods. To estimate this covariance, one estimates the covariance of non-averaged data and then, assuming independence of each trial, the estimated covariance gets divided by the number of trials present in the average (Engemann et al., 2015). Although the interpolation step can slightly bias the estimation, the noise covariance can still be estimated from single trials denoised by *autoreject (local)* method. As in current pipelines, this procedure provides an integer number of remaining trials for covariance scaling as each trial has the same weight in the average. This is in contrast to methods such as robust regression (Diedrichsen and Shadmehr, 2005) where a weighted average is computed with noisy trials down weighted.

Finally, a discussion for handling bad segments in M/EEG would not be complete without mentioning the quick and simple method that consists in declaring bad data segments as missing data. Practically put, the idea is to replace bad segments by the Not-a-Number (NaN) value, so that the trials can be averaged ignoring these NaNs (using for example the *nanmean* command). This approach is a simple alternative to sensor interpolation. Yet, this approach has some limitations. First, it does not give access to clean single trials with a common set of sensors. Second, the resulting average is obtained using different number of trials for each sensor. The noise variance reduction is not the same for all sensors. As a consequence, the noise covariance used for whitening the average data prior to inverse modeling cannot be obtained by dividing the covariance computed on raw data by an integer. Indeed, this integer number that matches the number of averaged trials is not the same for all sensors using this simple approach. In other words, this method is not readily compatible with common source localization methods.

6 Conclusion

In summary, we have presented a novel algorithm for automatic data-driven detection and repair of bad segments in single trial M/EEG data. We therefore termed it *autoreject*. We have compared our method to state-of-the-art methods on four different open datasets containing in total more than 200 subjects. Our validation suggests that *autoreject* performs at least as good as diverse alternatives and commonly used procedures while often performing considerably better. This is the consequence of the combination of a data-driven outlier-detection step combined with physics-driven channel repair where all parameters are calibrated using a cross-validation strategy robust to outliers. The insight about the necessity to tune parameters at the level of single sensors and for individual subjects was further consolidated by our analyses of threshold distributions. The empirical variability of optimal thresholds across datasets emphasizes the importance of statistical learning approaches and automatic model selection strategies for preprocessing M/EEG signals. While *autoreject* makes use of advanced statistical learning techniques such as Bayesian hyperparameter optimization, it is also grounded in the physics underlying the data generation. It is therefore not purely a black-box data-driven approach. It balances the trade-off between accuracy and interpretability. Indeed all *autoreject* parameters have a meaning from a user standpoint and the algorithmic decisions can be explained.

Supplemented by efficient diagnostic visualization routines, *autoreject* can be easily integrated in MEG/EEG analysis pipelines, including clinical ones where understanding algorithmic decisions is mandatory for tool adoption.

By offering an automatic and data-driven algorithmic solution to a task mostly so far done manually, *autoreject* reduces the cost of data inspection by experts. By allowing to repair data rather than removing it from the study, it allows saving data which are also costly to acquire. In addition, it removes the experts' bias which are due to specific training or prior experience, as well as some expectations about the data. It does so by defining a clear set of rules serving as inclusion criteria for M/EEG data, making results more easily reproducible and eventually limiting the risk of false discoveries. Furthermore, as data sharing across centers has become a common practice, *autoreject* it addresses the issue of heterogeneous acquisition setups. Indeed, each acquisition set-up has its intrinsic signal qualities, which means that preprocessing parameters can vary significantly between datasets. As opposed to alternative methods, *autoreject* automates the estimation of its parameters.

Finally, we want to emphasize that to the best of our knowledge, this study constitutes one of the first efforts in reanalysis of the HCP data for MEG using an entirely different stack of tools. The convergence between our method and the HCP MEG pipelines is encouraging and testifies to the success of the community-wide open science efforts aiming at reproducible research. We will therefore naturally make the code available online.

7 Acknowledgement

We thank Lionel Naccache for providing us with dramatic examples of artifact-ridden clinical EEG data which considerably stimulated the research presented in this study. The work was supported by the French National Research Agency (ANR-14-NEUC-0002-01) and the National Institutes of Health (R01 MH106174). Denis A. Engemann acknowledges support by the Amazon Webservices Research Grant awarded to him and the ERC StG 263584 awarded to Virginie van Wassenhove. We thank the Cognitive Neuroimaging Unit at Neurospin for fruitful discussions and feedback on this work. We further thank the MNE-Python developers and the MNE community for continuous collaborative interaction on basic development of the academic MNE software without which this study could not have been conducted.

References

- K. J. Gorgolewski, R. Poldrack, A practical guide for improving transparency and reproducibility in neuroimaging research, *bioRxiv* (2016) 039354.
- D. Van Essen, K. Ugurbil, E. Auerbach, D. Barch, T. Behrens, R. Bucholz, A. Chang, L. Chen, M. Corbetta, S. Curtiss, et al., The Human Connectome Project: a data acquisition perspective, *NeuroImage* 62 (2012) 2222–2231.
- L. J. Larson-Prior, R. Oostenveld, S. Della Penna, G. Michalareas, F. Prior, A. Babajani-Feremi, J.-M. Schoffelen, L. Marzetti, F. de Pasquale, F. Di Pompeo, et al., Adding dynamics to the Human Connectome Project with MEG, *NeuroImage* 80 (2013) 190–201.
- A. Goldberger, L. Amaral, L. Glass, J. Hausdorff, P. Ivanov, et al., Physiobank, physiotoolkit, and physionet components of a new research resource for complex physiologic signals, *Circulation* 101 (2000) e215–e220.

- G. Niso, C. Rogers, J. T. Moreau, L.-Y. Chen, C. Madjar, S. Das, E. Bock, F. Tadel, A. C. Evans, P. Jolicœur, et al., OMEGA: the open MEG archive, *NeuroImage* 124 (2016) 1182–1187.
- K. J. Gorgolewski, T. Auer, V. D. Calhoun, R. C. Craddock, S. Das, E. P. Duff, G. Flandin, S. S. Ghosh, T. Glatard, Y. O. Halchenko, et al., The Brain Imaging Data Structure: a standard for organizing and describing outputs of neuroimaging experiments, *bioRxiv* (2016) 034561.
- N. Bigdely-Shamlo, K. Kreutz-Delgado, K. Robbins, M. Miyakoshi, M. Westerfield, T. Bel-Bahar, C. Kothe, J. Hsi, S. Makeig, Hierarchical event descriptor (HED) tags for analysis of event-related EEG studies, in: *Global Conference on Signal and Information Processing (GlobalSIP)*, IEEE, pp. 1–4.
- W. Ollier, T. Sprosen, T. Peakman, UK Biobank: from concept to reality (2005).
- D. A. Engemann, A. Gramfort, Automated model selection in covariance estimation and spatial whitening of meg and eeg signals, *NeuroImage* 108 (2015) 328–342.
- D. Cruse, S. Chennu, C. Chatelle, T. A. Bekinschtein, D. Fernández-Espejo, J. D. Pickard, S. Laureys, A. M. Owen, Bedside detection of awareness in the vegetative state: a cohort study, *The Lancet* 378 (2012) 2088–2094.
- A. M. Goldfine, J. C. Bardin, Q. Noirhomme, J. J. Fins, N. D. Schiff, J. D. Victor, Reanalysis of bedside detection of awareness in the vegetative state: a cohort study., *Lancet* 381 (2013) 289.
- A. Basirat, S. Dehaene, G. Dehaene-Lambertz, A hierarchy of cortical responses to sequence violations in three-month-old infants, *Cognition* 132 (2014) 137–150.
- S. Taulu, M. Kajola, J. Simola, Suppression of interference and artifacts by the signal space separation method, *Brain Topography* 16 (2004) 269–275.
- F. Tadel, S. Baillet, J. C. Mosher, D. Pantazis, R. M. Leahy, Brainstorm: a user-friendly application for MEG/EEG analysis, *Computational intelligence and neuroscience* 2011 (2011) 8.
- A. Delorme, S. Makeig, EEGLAB: an open source toolbox for analysis of single-trial EEG dynamics including independent component analysis, *Journal of neuroscience methods* 134 (2004) 9–21.
- R. Oostenveld, P. Fries, E. Maris, J.-M. Schoffelen, FieldTrip: open source software for advanced analysis of MEG, EEG, and invasive electrophysiological data, *Computational intelligence and neuroscience* 2011 (2011).
- A. Gramfort, M. Luessi, E. Larson, D. Engemann, D. Strohmeier, et al., MEG and EEG data analysis with MNE-Python, *Front. Neurosci.* 7 (2013).
- V. Litvak, J. Mattout, S. Kiebel, C. Phillips, R. Henson, J. Kilner, G. Barnes, R. Oostenveld, J. Daunizeau, G. Flandin, et al., EEG and MEG data analysis in SPM8, *Computational intelligence and neuroscience* 2011 (2011).
- H. Nolan, R. Whelan, R. Reilly, FASTER: Fully Automated Statistical Thresholding for EEG artifact Rejection, *J. Neurosci. Methods* 192 (2010) 152–162.

- N. Bigdely-Shamlo, T. Mullen, C. Kothe, K.-M. Su, K. Robbins, The PREP pipeline: standardized preprocessing for large-scale EEG analysis, *Front. Neuroinform.* 9 (2015).
- A. Barachant, A. Andreiev, M. Congedo, The Riemannian Potato: an automatic and adaptive artifact detection method for online experiments using Riemannian geometry, in: *TOBI Workshop IV*, pp. 19–20.
- J. Diedrichsen, R. Shadmehr, Detecting and adjusting for artifacts in fMRI time series data, *NeuroImage* 27 (2005) 624–634.
- A. De Cheveigné, J. Simon, Sensor noise suppression, *J. of Neurosci. Methods* 168 (2008) 195–202.
- J. Snoek, H. Larochelle, R. P. Adams, Practical bayesian optimization of machine learning algorithms, in: *Advances in Neural Information Processing Systems*, pp. 2951–2959.
- J. S. Bergstra, R. Bardenet, Y. Bengio, B. Kégl, Algorithms for hyper-parameter optimization, in: *Advances in Neural Information Processing Systems*, pp. 2546–2554.
- D. Wakeman, R. Henson, A multi-subject, multi-modal human neuroimaging dataset, *Sci. Data* 2 (2015).
- G. Schalk, D. J. McFarland, T. Hinterberger, N. Birbaumer, J. R. Wolpaw, BCI2000: a general-purpose brain-computer interface (BCI) system, *IEEE Transactions on Biomedical Engineering* 51 (2004) 1034–1043.
- M. Jas, D. Engemann, F. Raimondo, Y. Bekhti, A. Gramfort, Automated rejection and repair of bad trials in MEG/EEG, in: *6th International Workshop on Pattern Recognition in Neuroimaging (PRNI)*.
- D. Leung, Cross-validation in nonparametric regression with outliers, *Ann. Stat.* (2005) 2291–2310.
- J. De Brabanter, K. Pelckmans, J. Suykens, J. Vandewalle, B. De Moor, Robust cross-validation score functions with application to weighted least squares support vector machine function estimation, *Technical Report*, K.U. Leuven, 2003.
- A. Gramfort, M. Luessi, E. Larson, D. Engemann, D. Strohmeier, et al., MNE software for processing MEG and EEG data, *NeuroImage* 86 (2014) 446 – 460.
- F. Perrin, J. Pernier, O. Bertrand, J. Echallier, Spherical splines for scalp potential and current density mapping, *Electroen. Clin. Neuro.* 72 (1989) 184–187.
- M. Hämäläinen, R. Ilmoniemi, Interpreting magnetic fields of the brain: minimum norm estimates, *Med. Biol. Eng. Comput.* 32 (1994) 35–42.
- C.-W. Hsu, C.-C. Chang, C.-J. Lin, et al., *A practical guide to support vector classification* (2003).
- A. Krizhevsky, I. Sutskever, G. E. Hinton, Imagenet classification with deep convolutional neural networks, in: *Advances in Neural Information Processing Systems*, pp. 1097–1105.
- C. E. Rasmussen, C. K. Williams, *Gaussian processes for machine learning*. 2006, The MIT Press, Cambridge, MA, USA 38 (2006) 715–719.

- F. Pedregosa, G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel, M. Blondel, P. Prettenhofer, R. Weiss, V. Dubourg, J. Vanderplas, A. Passos, D. Cournapeau, M. Brucher, M. Perrot, E. Duchesnay, Scikit-learn: Machine learning in Python, *Journal of Machine Learning Research* 12 (2011) 2825–2830.
- M. A. Fischler, R. C. Bolles, Random sample consensus: a paradigm for model fitting with applications to image analysis and automated cartography, *Communications of the ACM* 24 (1981) 381–395.
- R. Vigário, J. Sarela, V. Jousmiki, M. Hamalainen, E. Oja, Independent component approach to the analysis of EEG and MEG recordings, *IEEE Transactions on Biomedical Engineering* 47 (2000) 589–593.
- J. Dammers, M. Schiek, F. Boers, C. Silex, M. Zvyagintsev, U. Pietrzyk, K. Mathiak, Integration of amplitude and phase statistics for complete artifact removal in independent components of neuromagnetic recordings, *IEEE Transactions on Biomedical Engineering* 55 (2008) 2353–2362.
- M. Junghöfer, T. Elbert, D. M. Tucker, B. Rockstroh, Statistical control of artifacts in dense array EEG/MEG studies, *Psychophysiology* 37 (2000) 523–532.
- D. Engemann, F. Raimondo, J. King, M. Jas, A. Gramfort, S. Dehaene, L. Naccache, J. Sitt, Automated measurement and prediction of consciousness in vegetative state and minimally conscious patients., in: *Workshop on Statistics, Machine Learning and Neuroscience at the International Conference on Machine Learning (ICML)*, Lille.
- R. Vigário, Extraction of ocular artefacts from EEG using independent component analysis, *Electroen. Clin. Neuro.* 103 (1997) 395–404.
- M. Uusitalo, R. Ilmoniemi, Signal-space projection method for separating MEG or EEG into components, *Med. Biol. Eng. Comput.* 35 (1997) 135–140.
- K. Uutela, M. Hämäläinen, E. Somersalo, Visualization of magnetoencephalographic data using minimum current estimates, *NeuroImage* 10 (1999) 173–180.
- A. M. Dale, A. K. Liu, B. R. Fischl, R. L. Buckner, J. W. Belliveau, J. D. Lewine, E. Halgren, Dynamic statistical parametric mapping: combining fMRI and MEG for high-resolution imaging of cortical activity, *Neuron* 26 (2000) 55–67.
- R. D. Pascual-Marqui, et al., Standardized low-resolution brain electromagnetic tomography (sLORETA): technical details, *Methods Find Exp Clin Pharmacol* 24 (2002) 5–12.
- D. Engemann, F. Raimondo, J.-R. King, M. Jas, A. Gramfort, S. Dehaene, L. Naccache, J. Sitt, Automated measurement and prediction of consciousness in vegetative and minimally conscious patients, in: *ICML Workshop on Statistics, Machine Learning and Neuroscience (Stamline 2015)*.