

Comparison of beamformers for EEG source signal reconstruction

Yaqub Jonmohamadi^{a,b,*}, Govinda Poudel^{b,c,d}, Carrie Innes^{b,e,f}, Daniel Weiss^{g,h},
Rejko Krueger^{g,h,i}, Richard Jones^{a,b,e,f}

^a Department of Medicine, University of Otago, Christchurch, New Zealand

^b New Zealand Brain Research Institute, Christchurch, New Zealand

^c Monash Biomedical Imaging, Monash University, Melbourne, Australia

^d School of Psychological Sciences, Monash University, Clayton, Victoria, Australia

^e Department of Electrical and Computer Engineering, University of Canterbury, Christchurch, New Zealand

^f Department of Medical Physics and Bioengineering, Christchurch Hospital, Christchurch, New Zealand

^g German Center for Neurodegenerative Diseases (DZNE), Tübingen, Germany

^h Hertie-Institute for Clinical Brain Research, University of Tübingen, Germany

ⁱ Luxembourg Center for Systems Biomedicine, University of Luxembourg and Centre Hospitalier de Luxembourg, Luxembourg

ARTICLE INFO

Article history:

Received 21 June 2013

Received in revised form 23 July 2014

Accepted 23 July 2014

Keywords:

Beamformer

EEG

Gain

Region of interest (ROI)

Signal-to-noise power ratio (SNR)

Time-course reconstruction

ABSTRACT

Recently, several new beamformers have been introduced for reconstruction and localization of neural sources from EEG and MEG. Although studies have compared the accuracy of beamformers for localization of strong sources in the brain, a comparison of new and conventional beamformers for time-course reconstruction of a desired source has not been previously undertaken. In this study, 8 beamformers were examined with respect to several parameters, including variations in depth, orientation, magnitude, and frequency of the simulated source to determine their (i) effectiveness at time-course reconstruction of the sources, and (ii) stability of their performances with respect to the input changes. The spatial and directional pass-bands of the beamformers were estimated via simulated and real EEG sources to determine spatial resolution. White-noise spatial maps of the beamformers were calculated to show which beamformers have a location bias. Simulated EEG data were produced by projection via forward head modelling of simulated sources onto scalp electrodes, then superimposed on real background EEG. Real EEG was recorded from a patient with essential tremor and deep brain implanted electrodes. Gain – the ratio of SNR of the reconstructed time-course to the input SNR – was the primary measure of performance of the beamformers.

Overall, minimum-variance beamformers had higher Gains and superior spatial resolution to those of the minimum-norm beamformers, although their performance was more sensitive to changes in magnitude, depth, and frequency of the simulated source. White-noise spatial maps showed that several, but not all, beamformers have an undesirable location bias.

© 2014 Elsevier Ltd. All rights reserved.

1. Introduction

Electroencephalography (EEG) and magnetoencephalography (MEG) are noninvasive tools for functional brain imaging using scalp recording. Compared with other common tools for brain functional imaging such as functional magnetic resonance imaging (fMRI) and positron emission tomography (PET), which measure relatively slow changes in blood flow and metabolic activity

which are indirect markers of brain electrical activity, EEG and MEG measure brain electrical activity with millisecond temporal resolution. This advantage provides opportunities for studies of highly dynamic and transient neural activity. In recent years, brain source imaging and reconstruction from continuous and single-trial EEG/MEG data have received increased attention aimed at improving the understanding of rapidly changing brain dynamics [1–3] and using this for improved real-time brain monitoring, brain computer interface (BCI), and neurofeedback [4–6]. In contrast, EEG and MEG have poor spatial resolution relative to fMRI and PET. This is, in part, due to EEG and MEG mostly reflecting the electrical activity of the cortical grey matter, with deeper brain activities attenuated and contributing considerably less to the EEG/MEG signals. The beamformer provides a versatile form of spatial filtering,

* Corresponding author at: New Zealand Brain Research Institute, 66 Stewart Street, Christchurch 8011, New Zealand. Tel.: +64 33786093.

E-mail addresses: jonya247@student.otago.ac.nz, jacob.jonmo@yahoo.ca (Y. Jonmohamadi).

suitable for processing data from an array of sensors [7]. Beamformers were originally applied in array signal processing including sonar, radar, and seismic exploration [8]. The basic principle of beamformer design is to allow the neuronal signal of interest to pass through in a certain source location and orientation, called a pass-band, while suppressing noise or unwanted signal in other locations or orientations, called a stop-band [9]. A major limitation of beamformers is that they cannot properly reconstruct two spatially separate but temporally-correlated sources [9–11]; for example, they cancel each other when spatially far from each other or merge when they are spatially placed close to each other [9].

In recent years, new beamformers have been introduced for brain source localization and signal reconstruction from EEG and MEG [7,9,12,13]. The performances of these beamformers have mostly been evaluated in terms of accuracy for source localization of strong electric/magnetic signals, such as epileptogenic spikes [10,14], auditory evoked potentials [7,13,15], and median-nerve evoked potentials [9]. Aside from source localization, the application of beamformers to signal reconstruction of predefined regions of interest (ROI) in the brain is gaining increased attention in neuroimaging laboratories and is a common process in applications of EEG and MEG [16]. Examples of such ROIs are the motor cortex for BCI [17], intracerebral current flow for neurofeedback [16], or any region in the brain found to have consistent changes in activity via functional imaging techniques such as fMRI or PET. However, a comparison of several new and conventional beamformers for time-course reconstruction of a desired source has not been previously undertaken.

Beamformers applied to EEG or MEG fall into two categories: (A) scalar beamformers, which reconstruct the source time-course via a single output, and (B) vector beamformers which, reconstruct the source time-course in 3 orthogonal directions. For scalar beamformers, the orientation of the brain source can be estimated via techniques such as grid search [18,19], whereas vector beamformers do not require orientation of brain sources as they reconstruct the source time-series in 3 orthogonal time-courses. There are two methods for implementation of vector beamformers which are discussed in [20]. In the first method, the vector beamformer is a single beamformer with 3 orthogonal outputs, as applied in [10]. In the second implementation, the vector beamformer is made of 3 scalar beamformers in the 3 orthogonal directions as in [9,21].

In the current study, we investigated the performance of 8 beamformers: (1) minimum-variance (MV) (also known as distortionless minimum-variance [12]), (2) weight-normalized minimum-variance (WNMV) [12] (also known as Borgiotti–Kaplan [7,22]), (3) standardized minimum-variance (SMV) [12], (4) eigenspace extension of the minimum-variance (ESMV) [23], (5) higher-order covariance matrix of minimum-variance (HOC) [9], (6) generalized sidelobe canceller form of quiescent beamformer (GSC) [24], (7) standardized low-resolution electromagnetic tomography (sLORETA) [25], and (8) array-gain constraint minimum-norm with recursively updated Gram matrix (AGMN-RUG) [13]. *Gain*, defined as the ratio of the input signal-to-noise ratio (SNR) to that of reconstructed time-course SNR, was used to quantify the performance of the beamformers in the ROI, with respect to changes in source parameters of depth, orientation, magnitude, and frequency. Spatial and directional pass-bands were provided to show the spatial resolution of the beamformers. White-noise spatial maps of the beamformers were obtained by back-projection of the white noise to the source-space via beamformers to determine which beamformers have a location bias. For the most part, the scalar beamformers were applied to determine the performance of the beamformers with respect to changes in the input parameters and the spatial resolution.

Throughout this paper, plain italics indicate scalars, lower-case boldface italics indicate vectors, and upper-case boldface italics

indicate matrices. Subscript b refers to assumed location or orientation of the source and subscript d refers to actual location or orientation of the source. The Frobenius norm was used to obtain the norm of the matrices and vectors.

2. Beamformer algorithms

The reconstructed source time-series $\hat{s}(t, \mathbf{r}_b, \mathbf{q}_b)$ from the EEG for a scalar beamformer is

$$\hat{s}(t, \mathbf{r}_b, \mathbf{q}_b) = \mathbf{w}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{b}(t), \quad (1)$$

where $\mathbf{r}_b = [r_{bx}, r_{by}, r_{bz}]^T$ (mm) and $\mathbf{q}_b = [q_{bx}, q_{by}, q_{bz}]^T$ are the assumed source location and orientation respectively for calculation of the beamformer weight vector, $\|\mathbf{q}_b\| = 1$, $\mathbf{w}(\mathbf{r}_b, \mathbf{q}_b)$ is the weight vector for the scalar beamformer, and $\mathbf{b}(t) = [b_1(t), b_2(t), \dots, b_M(t)]^T$ is the measured EEG data on M electrodes at time t . Each beamformer has its own formulation of $\mathbf{w}(\mathbf{r}_b, \mathbf{q}_b)$.

For a vector beamformer the reconstructed time-course can be written as

$$\hat{\mathbf{s}}(t, \mathbf{r}_b) = \mathbf{W}^T(\mathbf{r}_b) \mathbf{b}(t). \quad (2)$$

In the above implementation, the vector beamformer is a single beamformer with 3 orthogonal outputs, as shown in [10]. A second implementation is shown in [9]:

$$\hat{s}_\mu(t, \mathbf{r}_b) = \mathbf{w}_\mu^T(\mathbf{r}_b) \mathbf{b}(t), \quad \mu = x, y, z. \quad (3)$$

Similar to [20], we call the first implementation a 3-D vector beamformer and the second implementation a 3-scalar vector beamformer.

Spatial filters can also be divided into two main families: minimum-variance and minimum-norm based spatial filters. In this study the MV, WNMV, SMV, HOC and ESMV beamformers belong to the minimum-variance family of spatial filters whereas GSC, sLORETA, and AGMN-RUG beamformers belong to the minimum-norm family of spatial filters. The minimum-variance spatial filters seek an adaptive solution for the minimization of the reconstructed source power. For the scalar beamformers, and without the loss of generality, this can be expressed as

$$\mathbf{w}(\mathbf{r}_b, \mathbf{q}_b) = \arg \min_{\mathbf{w}(\mathbf{r}_b, \mathbf{q}_b)} (\mathbf{w}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{C} \mathbf{w}(\mathbf{r}_b, \mathbf{q}_b)) \quad (4)$$

subject to

$$\mathbf{w}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{l}(\mathbf{r}_b, \mathbf{q}_b) = 1 \quad \text{for MV, ESMV, and HOC,}$$

$$\mathbf{w}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{w}(\mathbf{r}_b, \mathbf{q}_b) = 1 \quad \text{for WNMV,} \quad (5)$$

$$\mathbf{w}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{l}(\mathbf{r}_b, \mathbf{q}_b) = (\mathbf{l}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{C}^{-1} \mathbf{l}(\mathbf{r}_b, \mathbf{q}_b))^{1/2} \quad \text{for SMV,}$$

and

$$\mathbf{l}(\mathbf{r}_b, \mathbf{q}_b) = \mathbf{L}(\mathbf{r}_b) \mathbf{q}_b. \quad (6)$$

$\mathbf{L}(\mathbf{r}_b) = [\mathbf{l}_x(\mathbf{r}_b), \mathbf{l}_y(\mathbf{r}_b), \mathbf{l}_z(\mathbf{r}_b)]$ mm is $M \times 3$ lead-field matrix which gives the sensitivities of M EEG sensors for an assumed source location at $\mathbf{r}_b$ and $\mathbf{C}$ is the covariance matrix of EEG channels

$$\mathbf{C} = \langle \mathbf{b}(t) \mathbf{b}^T(t) \rangle, \quad (7)$$

where $\langle \dots \rangle$ is the ensemble average. Since the minimum-variance beamformers require the inverse of the covariance matrix $\mathbf{C}^{-1}$, a problem may arise when the $\mathbf{C}$ is not full rank. Brookes et al. [26] have suggested using a long window (i.e., as long as possible) of sensor data for calculation of $\mathbf{C}$, which helps $\mathbf{C}$ retain its full rank and obtain a better spatial resolution. However, using a long window of sensor data increases the risk of including the artefacts which happen frequently during recordings. Alternatively, when the use of a long window is not possible or the input SNR is high then $\mathbf{C}$ will not be a full rank matrix and the regularized inverse is suggested $(\mathbf{C} + \gamma \mathbf{I})^{-1}$ instead of $\mathbf{C}^{-1}$, where $\mathbf{I}$ is the unitary matrix,

$\gamma = 0.003\lambda_1$ is the regularization factor, and λ_1 is the largest eigenvalue of $\mathbf{C}$ [7]. The regularized inverse will also increase the *Gain* of the beamformer [18,27,28] but leads to higher interference from other sources close to the source of interest signal and reduces the spatial resolution; that is, there is a trade-off between spatial resolution and *Gain*. The effect of regularized inverse $\mathbf{C}$ on performance of the MV beamformer is well discussed in [26,29]. In this study we applied $\gamma = 0.001\lambda_1$ for regularization of the covariance matrix for all situations since $\mathbf{C}$ was not full rank in some cases.

The second family is that of minimum-norm in which the weight vector is derived by minimization of

$$\mathbf{w}(\mathbf{r}_b, \mathbf{q}_b) = \arg \min_{\mathbf{w}(\mathbf{r}_b, \mathbf{q}_b)} (\mathbf{w}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{G}_{\Omega} \mathbf{w}(\mathbf{r}_b, \mathbf{q}_b)) \quad (8)$$

subject to

$$\begin{aligned} \mathbf{w}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{l}(\mathbf{r}_b, \mathbf{q}_b) &= 1 \quad \text{for GSC (quiescent),} \\ \mathbf{w}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{G}_{\Omega} \mathbf{w}(\mathbf{r}_b, \mathbf{q}_b) &= 1 \quad \text{for sLORETA and AGMN – RUG,} \end{aligned} \quad (9)$$

where $\mathbf{G}_{\Omega}$ is the gram matrix

$$\mathbf{G}_{\Omega} = \sum_{\mathbf{r}_b \in \Omega} \mathbf{L}(\mathbf{r}_b) \mathbf{L}^T(\mathbf{r}_b) \quad (10)$$

and Ω is the ROI which can include several voxels dependent upon the size of ROI. The sLORETA and AGMN-RUG require the inverse of gram matrix. For inverse gram matrix the regularized inverse was used $(\mathbf{G}_{\Omega} + \gamma \mathbf{I})^{-1}$, where γ is the regularization parameter. Options for the regularization parameter have been rigorously investigated [30]. Typically, γ is set to the average of the variance of the sensor noise [31].

The normalized lead-field vector $\tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)$ is used in this study for all beamformers. The use of the normalized lead-field vector avoids the potential norm artefact of the lead-field vector [32], in which the norm of lead-field can change for different locations. For example in the case of spherical head model, the norm of lead-field becomes zero at the centre of the sphere and the weight vector for the MV beamformer become infinity.

$$\tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b) = \frac{\mathbf{l}(\mathbf{r}_b, \mathbf{q}_b)}{\|\mathbf{l}(\mathbf{r}_b, \mathbf{q}_b)\|}. \quad (11)$$

2.1. Minimum-variance beamformer

The MV beamformer is the best known beamformer in EEG and MEG applications [32]. Its weight vector $\mathbf{w}^T(\mathbf{r}_b, \mathbf{q}_b)$ for an assumed location and orientation of the source is

$$\mathbf{w}_{MV}(\mathbf{r}_b, \mathbf{q}_b) = \frac{\mathbf{C}^{-1} \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)}{\tilde{\mathbf{l}}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{C}^{-1} \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)}. \quad (12)$$

2.2. Weight-normalized minimum-variance beamformer

The WNMV beamformer was proposed for EEG and MEG applications in [7]. Since the constraint of WNMV is $\mathbf{w}_{WNMV}^T \times \tilde{\mathbf{l}} = \tau$ and $\mathbf{w}_{WNMV}^T \times \mathbf{w}_{WNMV} = 1$, therefore

$$\mathbf{w}_{WNMV}(\mathbf{r}_b, \mathbf{q}_b) = \tau \mathbf{w}_{MV}(\mathbf{r}_b, \mathbf{q}_b), \quad \tau = \frac{\tilde{\mathbf{l}}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{C}^{-1} \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)}{\sqrt{\tilde{\mathbf{l}}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{C}^{-2} \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)}}. \quad (13)$$

$$\mathbf{w}_{WNMV}(\mathbf{r}_b, \mathbf{q}_b) = \frac{\mathbf{C}^{-1} \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)}{\sqrt{\tilde{\mathbf{l}}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{C}^{-2} \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)}}. \quad (14)$$

The WNMV beamformer uses a normalized weight vector $\mathbf{w}_{WNMV}^T \times \mathbf{w}_{WNMV} = 1$ which ensures the signal reconstructed from any location has the same gain and, hence, avoids location bias, even when the non-normalized lead-field is used.

2.3. Standardized minimum-variance beamformer

The SMV beamformer was introduced in [12] and its vector version in [33]. The constraint for SMV is $\mathbf{w}_{WNMV}^T \times \tilde{\mathbf{l}} = \tau$ and

$$\tau = \sqrt{\tilde{\mathbf{l}}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{C}^{-1} \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)}, \quad (15)$$

$$\mathbf{w}_{SMV}(\mathbf{r}_b, \mathbf{q}_b) = \frac{\mathbf{C}^{-1} \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)}{\sqrt{\tilde{\mathbf{l}}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{C}^{-1} \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)}}. \quad (16)$$

In the case of formulation, the SMV beamformer is in fact the minimum-variance version of the well known sLORETA, i.e., substituting the $\mathbf{C}$ by $\mathbf{G}_{\Omega}$ yields sLORETA. The constraint of SMV beamformers ensures that the power of the reconstructed signal from any location in the brain is standardized with respect to $\sqrt{\tilde{\mathbf{l}}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{C}^{-1} \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)}$ which avoids the location bias. Thus, if there is only white noise on the sensors ($\mathbf{C} = \mathbf{I}$) then for all locations in the source space

$$\tau = \sqrt{\tilde{\mathbf{l}}^T(\mathbf{r}_b, \mathbf{q}_b) \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)} = 1 \quad (17)$$

and, therefore, SMV has a normalized white-noise spatial map and, hence, no location bias.

2.4. Eigen-space based minimum-variance beamformer

The ESMV beamformer [23] is based on eigen decomposition of the covariance matrix into noise and signal subspaces

$$\mathbf{C} = \mathbf{E}_S \mathbf{\Lambda}_S \mathbf{E}_S^T + \mathbf{E}_N \mathbf{\Lambda}_N \mathbf{E}_N^T, \quad (18)$$

where

$$\mathbf{\Lambda}_S = \text{diag}[\lambda_1, \lambda_2, \dots, \lambda_J] \quad \text{and} \quad \mathbf{\Lambda}_N = \text{diag}[\lambda_{J+1}, \lambda_{J+2}, \dots, \lambda_M], \quad (19)$$

where $\text{diag}[\dots]$ is the diagonal matrix with its elements being the eigenvalues of the signal space $\mathbf{\Lambda}_S$ and noise space $\mathbf{\Lambda}_N$, and the columns in $\mathbf{E}_S$ and $\mathbf{E}_N$ being the eigenvectors of $\mathbf{C}$. J is the number of the eigenvalues of the signal space. The weight vector for the ESMV beamformer is

$$\mathbf{w}_{ESMV}(\mathbf{r}_b, \mathbf{q}_b) = \mathbf{E}_S \mathbf{\Lambda}_S^{-1} \mathbf{w}_{MV}(\mathbf{r}_b, \mathbf{q}_b). \quad (20)$$

In Eq. (19) the eigenvalues from 1 to J belong to signal space and have greater values than σ^2 , where σ^2 is the variance of the background EEG. A difficulty associated with the application of ESMV is how to estimate J . At source locations, the lead-field vector is orthogonal to the noise subspace of the covariance matrix [34]

$$\tilde{\mathbf{l}}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{E}_N = 0, \quad (21)$$

but, in practice, the left side of Eq. (21) never becomes equal to zero [35]. In fact, J can be any number from 1 to M depending on the magnitude of the source. This is a limitation of the ESMV beamformer as it needs user information to define J . A work-around for this difficulty, suggested by [32], is to overestimate J (i.e., $J + \Delta J$), which gives a marginal drop in SNR compared with the correct estimation of J . However, the user may underestimate J and remove the desired signal from the signal space. As a result, the ESMV beamformer is more desirable for cases when the SNR is high and, therefore, the eigenvalue of the signal space is several times greater than that of the noise space and J will be a small number such as 1, 2 or 3.

2.5. Higher-order covariance matrix minimum-variance beamformer

The HOC beamformer was introduced in [9] and has a similar formula to that of the MV beamformer except for the use of higher-order terms (e.g., 2 or 3) in the covariance matrix

$$\mathbf{w}_{HOC}(\mathbf{r}_b, \mathbf{q}_b) = \frac{\mathbf{C}^{-n} \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)}{\tilde{\mathbf{l}}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{C}^{-n} \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)}, \quad n = 2 \text{ or } 3. \quad (22)$$

In our study $n = 3$. In [9] the HOC was compared with other beamformers, including MV and WNMV, and was shown to be superior in the case of source localization by neural activity index for strong sources, but no example was given of time-course reconstruction.

2.6. Generalized sidelobe canceller of quiescent beamformer

The GSC beamformer has two channels, the first being the forward channel which can be any type of beamformer and the second being the blocking (also called nulling) channel. In this paper, the GSC used in [24] was implemented. The forward channel for this GSC beamformer is the minimum-norm filter (also called quiescent beamformer)

$$\mathbf{w}_f(\mathbf{r}_b, \mathbf{q}_b) = \frac{\tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)}{\tilde{\mathbf{l}}^T(\mathbf{r}_b, \mathbf{q}_b) \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)}, \quad (23)$$

whereas the blocking channel includes a matrix $\mathbf{B}_n$ which is in the null space of the lead-field matrix and aims to pass the background signal and not the desired signal

$$\mathbf{B}_n = \text{null}[\tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)], \quad (24)$$

where $\text{null}[\dots]$ refers to null space. Because of this nulling channel, GSC beamformers are also called nulling beamformers. The $\hat{\mathbf{s}}(t, \mathbf{r}_b, \mathbf{q}_b)$ via GSC is obtained by

$$\hat{\mathbf{s}}(t, \mathbf{r}_b, \mathbf{q}_b) = \mathbf{w}_f^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{b}(t) - (\mathbf{b}^T(t) \mathbf{B}_n) \mathbf{w}_n(t), \quad (25)$$

where $\mathbf{w}_n(t)$ is the weight vector of an adaptive algorithm such as least-mean-square (LMS) or recursive-least-square (RLS) [24]. Although GSC has been called the generalized sidelobe canceller form of the minimum-variance beamformer [24], the GSC described in [24] belongs to the minimum-norm family of spatial filters as its forward channel $\mathbf{w}_f$ is a minimum-norm filter. Indeed it is possible to derive the generalized sidelobe canceller for any beamformer by adding the blocking channel parallel to the beamformer. The GSC described in [24] is one of the earliest beamformers used in EEG and MEG signal processing.

2.7. Standardized low-resolution electromagnetic tomography

SLORETA [25] is a form of spatial filter which has a similar weight vector to the SMV beamformer but uses the gram matrix instead of the covariance matrix. SLORETA is in fact the minimum-norm version of SMV beamformer. Hence, sLORETA is independent of the measured data in calculating the weight vector

$$\mathbf{w}_{SLORETA}(\mathbf{r}_b, \mathbf{q}_b) = \frac{\mathbf{G}_{\Omega}^{-1} \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)}{\sqrt{\tilde{\mathbf{l}}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{G}_{\Omega}^{-1} \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)}}. \quad (26)$$

2.8. Array-gain constraint minimum-norm with recursively updating gram matrix beamformer

In the AGCMN-RUG beamformer [13], the weight vector is

$$\mathbf{w}_{AGCMN-RUG}(t, \mathbf{r}_b, \mathbf{q}_b) = \frac{\bar{\mathbf{G}}_{\Omega}^{-1}(t) \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)}{\tilde{\mathbf{l}}^T(\mathbf{r}_b, \mathbf{q}_b) \bar{\mathbf{G}}_{\Omega}^{-1}(t) \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b)}, \quad (27)$$

where, for a single point ROI, the gram matrix is

$$\bar{\mathbf{G}}_{\Omega}(t) = \tilde{\mathbf{l}}(\mathbf{r}_b, \mathbf{q}_b) \hat{\mathbf{s}}^2(t, \mathbf{r}_b, \mathbf{q}_b) \tilde{\mathbf{l}}^T(\mathbf{r}_b, \mathbf{q}_b). \quad (28)$$

The AGCMN-RUG beamformer is similar to sLORETA, with the main difference being that in AGCMN-RUG the gram matrix is updated for each time sample and, hence, the weight vector also needs to be calculated for each time sample. In the original paper [13], an important advantage of the AGCMN-RUG beamformer was claimed to be that it does not require the covariance matrix of the measured signal. However, after providing the formulations, Greenblatt et al. [12] mentions that in practice this beamformer is influenced by measured noise and, consequently, a regularized version of the AGCMN-RUG was proposed as a solution which uses the covariance matrix of the measured signal to regularize the gram matrix. The regularized version was used in this study.

3. Methods

3.1. Forward problem

The simulated EEG data on M electrodes for N time samples is

$$\mathbf{b}(t) = \begin{cases} \mathbf{n}(t), & t = 0, 1, \dots, \frac{N}{2} - 1 \\ (\mathbf{L}(\mathbf{r}_d) \mathbf{q}) s(t) + \mathbf{n}(t), & t = \frac{N}{2}, \dots, N - 1 \end{cases} \quad (29)$$

where $\mathbf{L}(\mathbf{r}_d) = [\mathbf{L}_x(\mathbf{r}_d), \mathbf{L}_y(\mathbf{r}_d), \mathbf{L}_z(\mathbf{r}_d)]$ is $M \times 3$ lead-field matrix which shows the sensitivity of M EEG electrodes for a dipole located at $\mathbf{r}_d = [r_{dx}, r_{dy}, r_{dz}]^T$ (mm) in the head, $\mathbf{q} = \alpha \mathbf{q}_d$, is the dipole moment in A.m, $\mathbf{q}_d = [q_{dx}, q_{dy}, q_{dz}]^T$ is the normalized dipole orientation for which $\|\mathbf{q}_d\| = 1$, α is the dipole magnitude, $s(t)$ is the dipole time-series with unit magnitude, and $\mathbf{n}(t)$ is the additive background signal which is real EEG. In this study, a sinusoidal signal was used as the dipole time-series

$$s(t) = \sin(2\pi f t), \quad (30)$$

where f is the frequency of the simulated dipole in Hz.

Four parameters of the simulated dipole were considered for examination of the beamformers: location $\mathbf{r}_d$, orientation $\mathbf{q}_d$, magnitude α , and frequency f .

3.2. Simulated EEG data

The simulated EEG data comprised 6s of background EEG $\mathbf{n}(t)$ followed by 6s of sinusoidal signal projected on the scalp via forward head modelling and superimposed on background EEG. The background EEGs were obtained from 3 subjects in a resting state mode. The covariance matrix was measured over the 12s window.

The boundary element method (BEM) model of the head [36] from the average MNI-template brain, implemented in the Field-Trip toolbox [37], was used for the forward solution with 3 layers and a conductivity ratio of skull to soft tissue of 0.0125. Background EEGs were from 3 healthy subjects in the resting state mode, sampled at 250 Hz, and pass-band filtered at 1–45 Hz.

3.3. Coordinate system

The MNI coordinate system was used to describe the spatial location of the simulated source. The directions of the x, y, and z axes are shown in Fig. 1 and the centre of the coordinates [0, 0, 0] is at the anterior commissure and in line with the anterior/posterior commissural line. The international 10/10 system was used to define the location of the 64 EEG electrodes and the reference electrode was between the Cz and CPz electrodes.

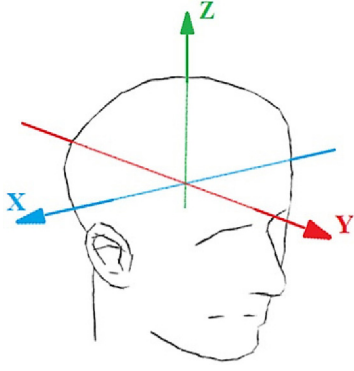


Fig. 1. The direction of the x , y , and z axes in the coordinate system used to describe the spatial location of the simulated source in the brain.

3.4. Performance measures

The performance of a beamformer was quantified by *Gain*, being the ratio of SNR of the reconstructed source time-series at the output of the beamformer, SNR_{out} , to SNR of the EEG data at the input to the beamformer, SNR_{in} ,

$$Gain = \frac{SNR_{out}}{SNR_{in}}. \quad (31)$$

SNR_{in} is the ratio of the power of the simulated source, Pin_{source} , divided by the power of the background signal, Pin_{noise} [10,16]. Pin_{source} was obtained via Frobenius norm of the $(L(r_d)q)s(t)$ in Eq. 29 for the 6 s in which the simulated source was active and Pin_{noise} was obtained via Frobenius norm of the $n(t)$ in Eq. 29 for the same 6 s window.

SNR_{out} is the ratio of the power of the reconstructed time-course $\hat{s}(t, r_b, q_b)$ at the frequency of the source, $P_{out_{source}}$, divided by the power of the reconstructed time-course at all frequencies except that of the simulated source

$$SNR_{out} = \frac{P_{out_{signal}}}{P_{out_{total}} - P_{out_{signal}}}. \quad (32)$$

To measure the power P , the FFT was applied to the 6 s window (6–12 s) of the reconstructed time-course of the source. The power of the desired source was measured at the frequency of the simulated source and power of the noise was the summation of power at all other frequencies.

The advantage of using *Gain* rather than SNR_{out} to determine the performance of the beamformer with respect to changes in input parameters is that *Gain* is independent of minor changes to the SNR_{in} , whereas SNR_{out} is dependent on SNR_{in} . Ideally, a beamformer's *Gain* would be independent of changes in source parameters such as depth, orientation, magnitude, and frequency. Hence, a major aim of this study was to compare the variance of beamformer's performances in relation to this ideal.

In the case of the ESMV beamformer, we found J via a search algorithm which tries all values of $J = 1, 2, \dots, M$, determines the *Gain* for each, and chooses the J with the highest *Gain*.

4. Computer simulations

Simulations were performed to (1) obtain and compare the spatial and directional pass-bands of the beamformers, (2) estimate the sensitivities of the beamformers to changes in the input parameters, and (3) obtain white-noise spatial maps of the scalar and vector beamformers.

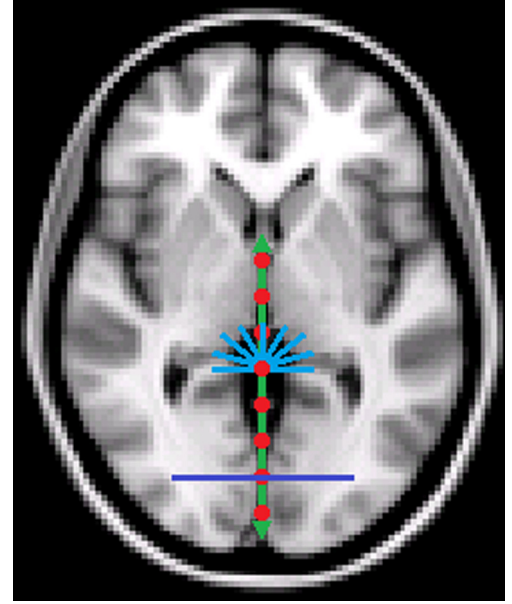


Fig. 2. This figure illustrates the different depths, orientations, and distances applied to measure the spatial pass-band of the beamformers. The red circles are the 8 depths (locations) of the simulated sources. As described in Section 4.2, for each of these locations, different orientations similar to the blue spokes were applied and for each location and orientation, different frequencies and magnitudes were assigned to the simulated source (details in Section 4.2). For the spatial pass-band (Section 4.1.1), the red circles show the locations of the simulated sources and the purple line shows the distances applied to obtain the spatial pass-band. For the directional pass-band (Section 4.1.2), the red circles are the locations where the simulated source was placed and the blue spokes are the orientations assigned to the beamformer while the orientation of the simulated source was the green line. (For interpretation of the references to colour in this figure legend, the reader is referred to the web version of this article.)

4.1. Pass-bands

Each scalar beamformer has a spatial and directional (orientation) pass-band. The spatial pass-band shows how effective the beamformer is at spatially allowing the desired source to pass while meantime suppressing sources nearby. Similarly, the directional pass-band shows how effective the beamformer is at allowing the desired source to pass at a desired angle while suppressing the activity at other angles. Thus, the pass-bands are measures of the spatial resolution of a beamformer, which is important in source imaging applications. An ideal beamformer would have pass-bands similar to the delta function. In this section the aim is to obtain the spatial and directional pass-bands of the beamformers.

4.1.1. Spatial pass-band

The spatial pass-band (also called 'beam response' [32]) can be estimated by measuring the *Gain* of beamformers at different distances from the simulated source. The highest *Gain* should be for zero distance. The spatial pass-band may also provide information on the localization bias of the beamformers; i.e., if the highest *Gain* is obtained at distances than zero, the beamformer has a localization bias. The simulated data were produced by placing a simulated source in the brain and running the beamformers at different distances from the source: -46 mm to $+46$ mm every 2 mm, as shown by purple line in Fig. 2. That is, for every simulated source there were 47 sample distances at which beamformers were placed. This process was applied for 8 depths in the head (red circles in Fig. 2), and 6 magnitudes of the simulated source, so as to cover both deep and shallow and weak and strong sources. This process was applied to the EEG background of 3 subjects. Therefore, the spatial pass-band for each beamformer comprises

Gains of 47 (sample distances) $\times$ 8 (depths) $\times$ 6 (SNRin) $\times$ 3 (EEG background) = 6768 conditions. In particular, the *Gain* at each sample distance comprises the mean and confidence interval of 144 iterations. The 8 depths are in steps of 10 mm from [0, –80, 0] mm to [0, 0, 0] mm. The 6 SNRins were 0.12, 0.35, 1.00, 3.45, 8.80, and 35.30. The locations of the simulated sources are shown in Fig. 2 by red circles and the purple line is an example of where the beamformers were run for one of the sample locations. The orientations of the simulated sources and beamformers were [0, 1, 0] in all cases.

4.1.2. Directional pass-band

The directional pass-band of beamformers can be estimated by measuring *Gain* of the beamformers in different orientations relative to the actual source orientation, while keeping the assumed location of the beamformer the same as for the simulated source. Ideally, the beamformer directional pass-band should be maximal when the angle between the beamformer and the source orientation is zero. The simulated data were produced by applying angles to the beamformer from –90 deg to +90 with respect to the actual source orientation. The directional pass-band for each beamformer comprises *Gain* of 47 (sample angles) $\times$ 8 (depths) $\times$ 6 (SNRin) $\times$ 3 (EEG background) = 6768 conditions. In particular, the *Gain* at each sample angle comprises the mean and confidence interval of 144 iterations. For each sample, the beamformer was placed at the same location as that of the simulated source. That is, 8 depth samples ([0, –80, 0] mm to [0, 0, 0] mm) with all orientations of simulated sources at [0, 1, 0]. The –90 deg corresponds to [–1, 0, 0] and +90 deg to [1, 0, 0]. In Fig. 2, the sample angles of the beamformers are shown as blue spokes, whereas the actual orientation of the source was [0, 1, 0], i.e., parallel to the green line. As for spatial pass-bands, the 6 SNRins were 0.12, 0.35, 1.00, 3.45, 8.80, and 35.30.

4.2. Effect of source parameters on the performance of beamformers

Here the aim was to determine the stability of the beamformer performance with respect to changes in the input parameters of depth, magnitude, orientation, and frequency. Ideally, *Gain* should remain constant with changes in any input parameters. A simulation was run which produced sources at different depths, magnitudes, orientations, and frequencies. The depths were from [0, –80, 0] mm to [0, 0, 0] mm (moving in y direction every 10 mm and shown by red circles in Fig. 2), the orientations were from [1, 0, 0] to [–1, 0, 0] (corresponding to 180 deg rotating around the z direction in 11 samples, shown as blue spokes in Fig. 2), SNRins 0.12, 0.35, 1.00, 3.45, 8.80, and 35.30, and the frequencies were 2, 5, 10, 20, and 50 Hz. All of the simulations were applied for the EEG backgrounds from 3 subjects. Therefore, there were 8 (depths) $\times$ 11 (orientations) $\times$ 6 (magnitudes) $\times$ 5 (frequencies) $\times$ 3 (EEG backgrounds) = 7920 iterations (*Gain* samples). In order to measure changes of *Gain* due to each of these 4 parameters, averaging was applied to other parameters; e.g., to obtain changes in performance with changes in SNRin, as for each of the 6 SNRin samples there were 8 (depths) $\times$ 11 (orientations) $\times$ 5 (frequencies) $\times$ 3 (EEG backgrounds) = 1320 iterations (*Gain* samples). Therefore, the mean and confidence interval of 1320 *Gain* samples were calculated and plotted. Similar processes were applied for each of the 4 parameters to show the sensitivity of beamformer performance to changes in each of these parameters. In all cases, the beamformers were given the same location and orientation as the simulated source.

4.3. Real EEG from deep brain stimulation

EEG was recorded from a patient with essential tremor (ET) treated with deep brain stimulation (DBS). Subject ET1 gave written informed consent. For DBS, bilateral quadripolar microelectrodes

were implanted stereotactically for ET treatment into the ventral intermedial thalami.

Each DBS electrode had 4 contacts, spaced 1.5 mm apart (lead model 3389, Medtronic, Meerbusch, Germany). A 64-channel 10–20 standard was used for EEG recording (BrainVision recorder, MES Electronics, Munich, Germany) and EEG was sampled at 5000 Hz and low-pass filtered at 1000 Hz.

A recording of 10 s of resting state EEG was made while the right stimulator (bipolar: contacts 0 and 1) was activated, at 10 Hz with pulses of 300 μ s width and amplitude of 1.0 V_{peak}. To estimate SNRin, independent component analysis (ICA) (infomax [38]) implemented in EEGLAB toolbox [39] was used to remove the DBS component (EEG_{DBS}) from other components of the EEG (EEG_{noise}). The FFT was applied to the DBS component of the 10 s epoch and the power of the DBS component (P_{DBS}) was defined as the summation of the power at the main frequency (10 Hz) and its harmonics. The remaining power was considered to be noise power P_{noise1} . This led to an SNR of the DBS component. As ICA does not retain the actual magnitude of the components, P_{DBS} and P_{noise1} , measured via FFT of the DBS component needed to be corrected (P'_{DBS} , P'_{noise1}) via the Frobenius norms of EEG_{DBS} scaled by the SNR of the DBS component. The power of the noise on the EEG_{noise}, P_{noise2} , was also determined via Frobenius norm of EEG_{noise}. From this,

$$SNRin = \frac{P'_{DBS}}{P'_{noise1} + P_{noise2}}, \quad (33)$$

and was found to be ≈ 0.065 .

As the DBS electrodes were implanted, it was not possible to change parameters such as orientation and location and, hence, only the spatial and directional pass-bands of the beamformers could be measured. Dipole fitting post-ICA [40] was used to estimate the anatomical location of the DBS component from scalp EEG as [10, –12, 0] mm, which is in the right thalamus where the DBS electrode indeed was implanted. In addition, we applied our method [41,42], which applies ICA in the source-space (on voxels) post-beamforming rather than sensor-space; a similar location was found ([6, –16, 2] mm). Based on the ICA, the normalized DBS orientation was estimated as [0.30, 0.55, 0.78], which converts to spherical coordinates of $[r, \theta, \Phi] = [1, 39, 28]$ deg. To obtain the spatial pass-band of the beamformers, while $y = -14$ mm and $z = 1$ mm, the beamformers were run for x from –50 mm to 60 mm in steps of 2 mm. To obtain the directional pass-band of the beamformers, the θ of the beamformers was varied from –90 to +90 deg relative to the DBS orientation, while Φ was kept constant at 28 deg.

Fig. 3(a) shows 0.5 s of the EEG of ET1, (b) is the time-course of the DBS component, (c) is the FFT of the DBS component, (d) is the topography of the DBS component (left) and the lead-field pattern of a simulated source in the right thalamus and the same orientation of the DBS component, and (e) is the source-imaging for the DBS component via post-beamforming ICA as described in [42].

4.4. White-noise spatial map

A beamformer may have a location bias, in which reconstructed time-courses from different parts of the brain have different magnitudes from identical sources. This is due to some beamformers having a non-uniform noise power map; this is shown in [10] for the vector MV beamformer. The reconstructed time-series for location $\mathbf{r}_b$ and orientation $\mathbf{q}_b$ at time t (equation 1) can be rewritten as

$$\begin{aligned} \hat{\mathbf{s}}(t, \mathbf{r}_b, \mathbf{b}_b) &= \mathbf{w}^T(\mathbf{r}_b, \mathbf{q}_b) \mathbf{b}(t) \\ &= \frac{\|\mathbf{w}(\mathbf{r}_b, \mathbf{q}_b)\| \mathbf{w}^T(\mathbf{r}_b, \mathbf{q}_b)}{\|\mathbf{w}(\mathbf{r}_b, \mathbf{q}_b)\|} \frac{\|\mathbf{b}(t)\|}{\|\mathbf{b}(t)\|} \mathbf{b}(t), \end{aligned} \quad (34)$$

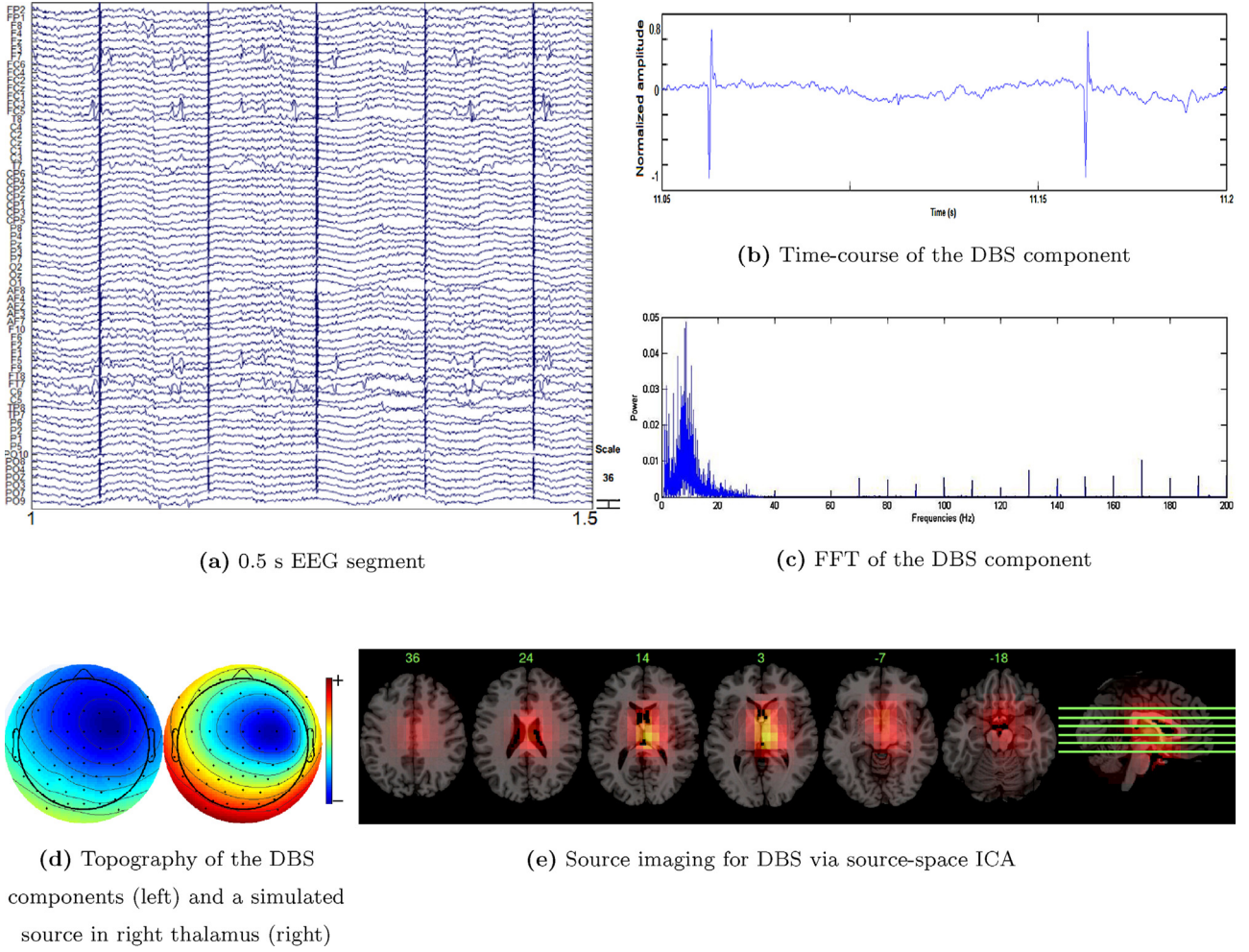


Fig. 3. (a) 0.5 s of EEG from subject ET1, with stimulation of the right DBS electrode placed in the right thalamus. The DBS comprised pulses at 10 Hz with a pulse width of 300 μ s. (b) DBS component separated by ICA from the recorded EEG. (c) FFT of the DBS component; power of DBS component was measured at 10 Hz and its harmonics up to $\sim$ 1 kHz (low-pass filter). (d) Topography of the DBS component identified by scalp ICA (left map) and topography of a simulated source placed in the right thalamus (right map). (e) Source imaging via source-space ICA [42] for the DBS EEG, with [6, –16, 2] mm as the focal point for the DBS.

where $\|\mathbf{w}(\mathbf{r}_b, \mathbf{q}_b)\|$ and $\|\mathbf{b}(t)\|$ are the norms of the beamformer weight vector and sensor signals respectively. Therefore, the power of the reconstructed time-series at location $\mathbf{r}_b$ depends on both the norm of the beamformer weight vector and the norm of the recorded neural activity on the sensors. If the norm of the beamformer weight vector changes due to the location, the power of the reconstructed time-series will also change, hence introducing location bias. This is why Van Veen et al. [10] proposed the neural activity index (NAI) (rather than power) measurement at each $\mathbf{r}_b$ as this normalizes the noise power map and reflects only changes in neural power. However, the NAI also rectifies and merges the orthogonal time-courses of the vector beamformer whereas the user may be interested in the actual time-courses of the sources of interest. Therefore, a beamformer which inherently has no location bias is of greater value when the primary aim is of time-course reconstruction of sources of interest. The WNMV beamformer is one of the beamformers which does not have a location bias due its weight vector being normalized such that for any $\mathbf{r}_b$ the norm of the beamformer weight vector is 1. This is why WNMV is also called the unit-noise-gain minimum-variance beamformer [32].

In source localization applications, location bias can result in both false and missed detections of actual sources [10]. Location bias can also be misleading when the user is interested in the

time-courses of activity from certain areas of interest as some locations can have magnitudes of higher variance due to location bias.

A good approach for identification of beamformers with location bias is to generate *white-noise spatial maps*, as applied in [10] for the vector MV beamformer. If the only signal on the sensors is white noise, $\|\mathbf{b}(t)\|$ will remain constant for all t and obtaining the power of the reconstructed time-series for different locations will reflect the norm of the weight vector for those locations.

We obtained white-noise spatial maps by applying beamformers on a 3-D scanning grid. Similar to [10], the only activity on the sensors was uncorrelated white noise (10 s). The beamformers were run on a 3-D scanning grid of 2041×10 mm³ voxels covering the whole brain. The power of each voxel was measured via $\sqrt{\langle \hat{s}^2(t, \mathbf{r}_d, \mathbf{q}_d) \rangle}$, then projected back onto the scanning grid to obtain white-noise spatial maps of the beamformers.

There are two implementations of vector beamformers: (a) vector beamformers with 3 orthogonal outputs (3-D vector beamformer as applied in [10]) and (b) vector beamformers which is made of 3 orthogonal scalar beamformers in each of the orthogonal directions (3-scalar vector beamformer as applied in [9]). In the case of vector MV, the 3-D vector beamformer can be written as

$$\mathbf{W}(\mathbf{r}_b) = \frac{\mathbf{C}^{-1}\tilde{\mathbf{L}}(\mathbf{r}_b)}{\tilde{\mathbf{L}}^T(\mathbf{r}_b)\mathbf{C}^{-1}\tilde{\mathbf{L}}(\mathbf{r}_b)} \quad (35)$$

whereas the 3-scalar vector beamformer is written as:

$$\mathbf{w}_\mu(\mathbf{r}_b) = \frac{\mathbf{C}^{-1}\tilde{\mathbf{l}}_\mu(\mathbf{r}_b)}{\tilde{\mathbf{l}}_\mu^T(\mathbf{r}_b)\mathbf{C}^{-1}\tilde{\mathbf{l}}_\mu(\mathbf{r}_b)}, \quad \mu = x, y, z. \quad (36)$$

The latter implementation was proposed by [9] as a means of improving source imaging via the neural activity index (NAI). Johnson et al. [20] have highlighted differences between these two implementations both for time-course reconstruction and localization via the MV beamformer. Similarly, the white-noise spatial maps can be different between two implementation. For example, if the signal on the sensors is white noise, $\mathbf{C} = \mathbf{I}$ and the norm of the weight vectors for the 3-D implementation is

$$\|\mathbf{W}(\mathbf{r}_b)\| = \frac{\|\tilde{\mathbf{l}}(\mathbf{r}_b)\|}{\|\tilde{\mathbf{l}}^T(\mathbf{r}_b)\tilde{\mathbf{l}}(\mathbf{r}_b)\|} = \|\tilde{\mathbf{l}}^{-1}(\mathbf{r}_b)\| \geq \frac{1}{\|\tilde{\mathbf{l}}(\mathbf{r}_b)\|} = \frac{1}{\sqrt{3}}. \quad (37)$$

In the above equation, the denominator $\|\tilde{\mathbf{l}}^T(\mathbf{r}_b)\tilde{\mathbf{l}}(\mathbf{r}_b)\|$ is not constant as it changes for each $\mathbf{r}_b$, since

$$\tilde{\mathbf{l}}^T(\mathbf{r}_b)\tilde{\mathbf{l}}(\mathbf{r}_b) = \begin{pmatrix} 1 & ? & ? \\ ? & 1 & ? \\ ? & ? & 1 \end{pmatrix}. \quad (38)$$

Therefore, the 3-D implementation of the MV beamformer does not have a uniform white-noise spatial map even when the normalized lead-field is used. For the 3-scalar vector implementation the norm of the weight vectors when only white noise is applied is

$$\|\mathbf{w}_\mu(\mathbf{r}_b)\| = \frac{\|\tilde{\mathbf{l}}_\mu(\mathbf{r}_b)\|}{\|\tilde{\mathbf{l}}_\mu^T(\mathbf{r}_b)\tilde{\mathbf{l}}_\mu(\mathbf{r}_b)\|} = 1, \quad \mu = x, y, z. \quad (39)$$

Therefore, the norm of the weight vectors for every $\mathbf{r}_b$ is 1 and the 3-scalar implementation of the vector MV beamformer has a uniform white-noise spatial map only when the normalized lead-field is used. The normalized lead-field is applied in this study, otherwise Eq. 39 is not correct and both implementations of the vector MV beamformer have a location bias. Eq. 39 also indicates that all scalar beamformers from the minimum-variance family have a uniform white-noise spatial map when the normalized lead-field is used.

We have calculated white-noise spatial maps for both vector implementations for all beamformers in our study to determine which implementations have or do not have location bias. The ESMV and HOC beamformers have the same structure as the MV beamformer. For the vector beamformers, the power of the time-courses at each location was measured via

$$\langle \|\hat{\mathbf{s}}(t, \mathbf{r}_b)\| \rangle = \sqrt{\langle \hat{\mathbf{s}}_x^2(t, \mathbf{r}_b) + \hat{\mathbf{s}}_y^2(t, \mathbf{r}_b) + \hat{\mathbf{s}}_z^2(t, \mathbf{r}_b) \rangle}. \quad (40)$$

5. Results

A qualitative summary of beamformer comparisons is given in Table 1. Quantitative results are in Figs 4–11. Gain characteristics of the MV, WNMV, SMV beamformers are near identical and,

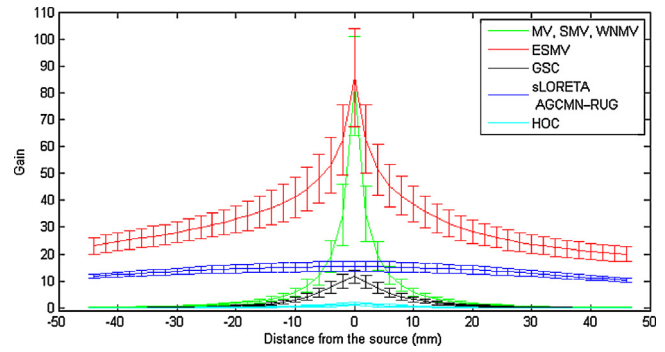


Fig. 4. Spatial pass-bands of the beamformers, with distance from source changed every 2 mm from –46 mm to +46 mm parallel to the x vector, as shown in Fig. 2 as a purple line. The vertical bars are mean $\pm$ 95% confidence interval.

hence, they are shown as single green lines in the graphs. Similarly, sLORETA and AGCMN-RUG were near identical in all simulations and are shown as single blue lines.

5.1. Pass-band

The spatial pass-bands of the beamformers are shown in Fig. 4. The MV, WNMV, and SMV beamformers have much narrower spatial pass-bands than the other beamformers. The full-width half-maximums (FWHM) for MV, WNMV, and SMV were 3.5 mm, ESMV 18 mm, GSC 12 mm, and sLORETA and AGCMN-RUG >88 mm. The HOC beamformer performed so poorly that it was not possible to estimate its FWHM. At 0 mm distance, the Gain of ESMV was slightly higher than MV, WNMV, and SMV. The maximum Gain for sLORETA and AGCMN-RUG was at –4 mm (4 mm in the –x direction), indicating a localization bias.

The directional pass-bands are shown in Fig. 5. The FWHMs for MV, WNMV, and SMV were 8 deg, ESMV 14 deg, GSC 128 deg, and sLORETA and AGCMN-RUG 62 deg. Again, the HOC beamformer had such a poor performance that it was not possible to estimate its FWHM. All of the beamformers had maximum Gain at 0 deg except for sLORETA and AGCMN-RUG which had maximums at 5 deg.

The spatial pass-band of the GSC is far narrower than its directional pass-band, whereas the opposite is the case for sLORETA and AGCMN-RUG. In the case of minimum-variance beamformers, both spatial and directional pass-bands were narrower compared to minimum-norm beamformers.

5.2. Source depth

The effect of depth of source on the Gain of the beamformers is shown in Fig. 6. Depth was defined as the distance of the simulated source, changing in the y direction from a scalp point at [0, –115, 0] mm. The Gain of minimum-variance beamformers (MV, WNMV, SMV, and ESMV) dropped considerably with increasing depths; e.g., at a depth of 95 mm, their Gains had decreased to ~13% of their

Table 1
Summary of beamformer comparisons.*

	Beamformer	Depth sensitivity	Magnitude sensitivity	Orientation sensitivity	Frequency sensitivity	Spatial resolution	Location bias
Minimum-variance	MV	High	High	Minimal	High	Highest	3-D vector version only
	WNMV	High	High	Minimal	High	Highest	None
	SMV	High	High	Minimal	High	Highest	None
	ESMV	High	High	Minimal	High	Intermediate	3-D vector version only
Minimum-norm	GSC	Medium	High	Medium	High	Intermediate	3-D vector version only
	sLORETA	Medium	Minimal	High	Minimal	Lowest	Scalar and vector versions
	AGCMN-RUG	Medium	Minimal	High	Minimal	Lowest	Scalar and vector versions

* ‘High’ means a reduction in Gain of > 50% of the maximum Gain; ‘Medium’ means a reduction between 50 and 20%; ‘Minimal’ means a reduction of < 20%.

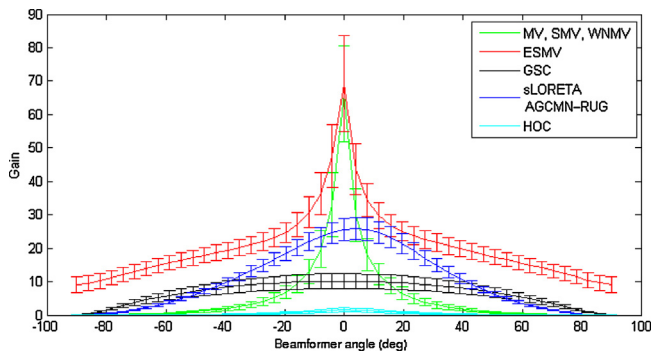


Fig. 5. Directional pass-bands of the beamformers. The angle of the beamformer was defined with respect to orientation of the simulated source which was parallel to the y vector. The angle of the beamformer is shown in Fig. 2 with blue spokes, whereas the angle of the simulated source was parallel to the green line. Therefore, the orientation of the source was $[0, 1, 0]$ and the orientation of the beamformer changed from $[1, 0, 0]$ to $[-1, 0, 0]$, corresponding to 180 deg around the z vector, every ~ 3.8 deg. The vertical bars are mean $\pm 95\%$ confidence interval. (For interpretation of the references to colour in this figure legend, the reader is referred to the web version of this article.)

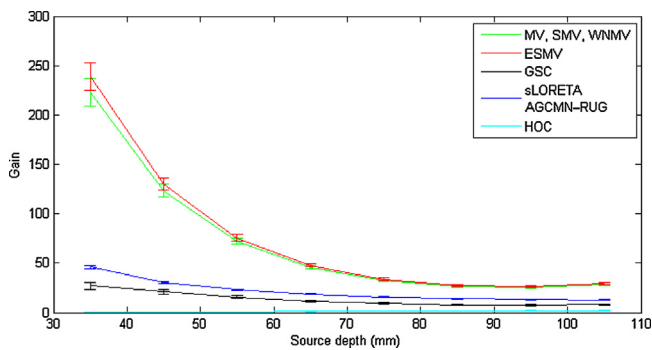


Fig. 6. Effect of depth of source on performance of the beamformers. The simulated source was placed at 8 depths, shown with red circles in Fig. 2. Depth was defined as the distance of the simulated source (which changes in y direction) to a scalp point at $[0, -115, 0]$ mm. The vertical bars are mean $\pm 95\%$ confidence interval. (For interpretation of the references to colour in this figure legend, the reader is referred to the web version of this article.)

Gains at 35 mm. Despite this, the Gains of minimum-variance beamformers were still higher than the other beamformers. An increase in the depth of the simulated source also caused a drop in Gain on the other beamformers, albeit to a smaller extent; e.g., at 95 mm the Gain of sLORETA and AGCMN-RUG had dropped to $\sim 34\%$ of the Gain at 35 mm and the Gain of GSC had dropped to $\sim 40\%$.

5.3. Source magnitude

The effect of magnitude of the source (SNR_{in}) on the Gain of the beamformers is shown in Fig. 7. The Gain of MV, WNMV, SMV, and ESMV dropped considerably with increasing SNR_{in} ; e.g., as SNR_{in} increased from 0.12 to 38.00 the Gain decreased by $\sim 81\%$ from 130 to 25 (~ 35 in the case of ESMV). GSC was the most affected by increase in magnitude of the sources, with a drop in Gain of $\sim 92\%$. sLORETA and AGCMN-RUG were linear in the case of changes in the magnitude of the source and retained a constant Gain for all SNR_{in} values.

5.4. Source orientation

The effect of different orientations on the performance of the beamformers is shown in Fig. 8. sLORETA and AGCMN-RUG were the most affected with $\sim 50\%$ fluctuation in Gain. GSC had a near

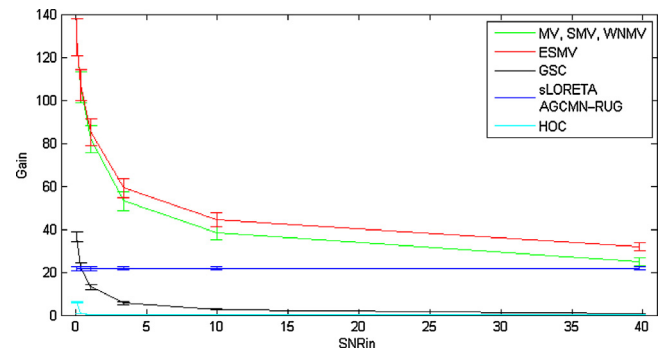


Fig. 7. Effect of varying SNR_{in} (magnitude of the source) on performance of the beamformers. $SNR_{in} = 0.12$ to 38.00 (6 samples). The vertical bars are mean $\pm 95\%$ confidence interval.

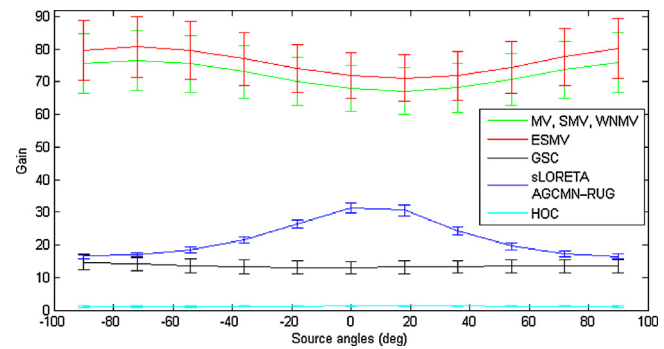


Fig. 8. Effect of varying orientation of the simulated source on performance of the beamformers. The different orientations (11 samples) for the source are shown in blue spokes in Fig. 2, corresponding to 180 deg around the z vector. The vertical bars are mean $\pm 95\%$ confidence interval. (For interpretation of the references to colour in this figure legend, the reader is referred to the web version of this article.)

zero change and MV, WNMV, SMV, and ESMV had relatively small changes in Gain of $\sim 10\%$.

5.5. Source frequency

The effect of different source frequencies on performance of the beamformers is shown in Fig. 9. MV, WNMV, SMV, ESMV, and GSC mostly increased in Gains with increased frequency of the simulated source; e.g., the Gains of MV WNMV, SMV, and ESMV at lowest frequency (2 Hz) were $\sim 40\%$ of their respective maximum Gains at highest frequency (50 Hz). Similarly, GSC at 2 Hz was only 10% of its maximum at 50 Hz. The Gains of sLORETA and AGCMN-RUG were independent of the frequency of the source.

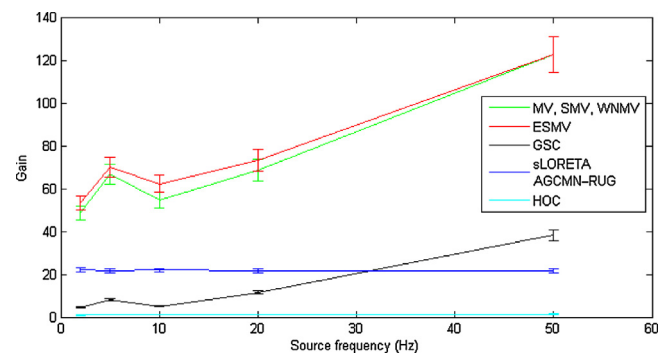


Fig. 9. Effect of varying the frequency of simulated source on performance of the beamformers. Five frequencies (2, 5, 10, 20, and 50 Hz) were assigned to the simulated source. The vertical bars are mean $\pm 95\%$ confidence interval.

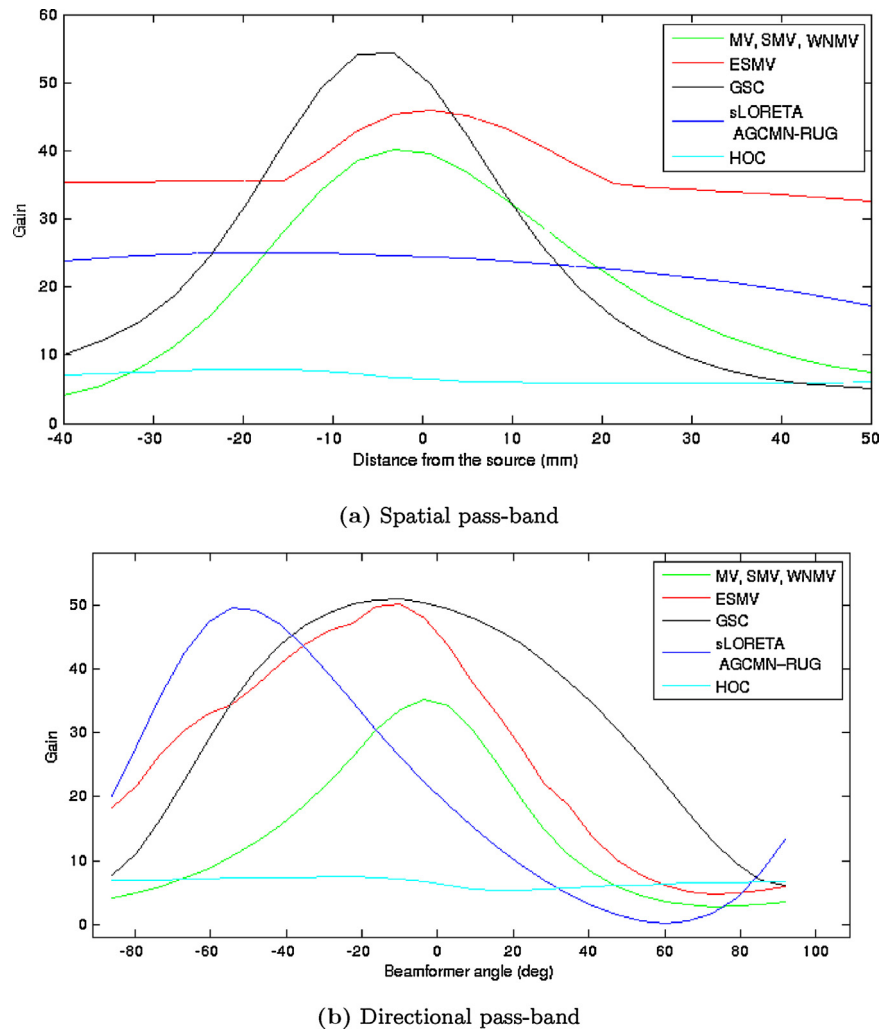


Fig. 10. Spatial and directional pass-bands of the beamformers for the DBS experiment. (a) Spatial pass-bands: distance from source was defined as the distance of the beamformer in the x direction from a point [8, 14, 1] mm where it was assumed the DBS electrode was located. The x of the beamformer was changing from -50 mm to $+60$ mm every 2 mm. The normalized DBS orientation was estimated to be [0.30, 0.55, 0.78] or in spherical coordinates, $[r, \theta, \Phi] = [1, 39, 28]$. (b) Directional pass-bands: θ was varied from -90 to $+90$ deg of angle difference in respect to the DBS orientation of θ , whereas Φ was kept constant at 28 deg.

5.6. Real EEG (DBS)

The spatial and directional pass-bands of the beamformers for EEG containing DBS signals are shown in Fig. 10(a) and (b) respectively. For the spatial pass-band, the FWHM of MV, WNMV, and SMV was 42 mm, which is slightly higher than the 34 mm for GSC. ESMV, sLORETA, and AGCMN-RUG had FWHMs more than 90 mm. For the directional pass-band, the FWHMs of MV, WNMV, and SMV were all 54 deg, compared with 106 deg for GSC, 94 deg for ESMV, and 72 deg for sLORETA and AGCMN-RUG.

There are similarities between the DBS spatial and directional pass-bands to those seen in the simulations, i.e., GSC had a narrower spatial than directional pass-band and, similarly, sLORETA and AGCMN-RUG had narrower directional than spatial pass-bands. Also, as for simulations, sLORETA, and AGCMN-RUG had maximum Gains different from those of the actual location or orientation of the DBS source.

There were, however, two main differences between the simulations and DBS results: (1) the spatial and directional pass-bands of the beamformers were considerably wider than the simulations, and (2) GSC had a higher Gain than MV, WNMV, and SMV, whereas in the simulations the GSC had an overall maximum Gain several times smaller than that of MV, WNMV, and SMV. The first difference

could be because the DBS electrodes are placed in a deep part of the brain and deeper sources tend to result in a more smearing of EEG topography which degrades the spatial resolution of the beamformers. Fig. 3(d) also shows differences between DBS topography and head modelling. The second difference could be due to the (1) depth effect which is also shown in Fig. 6 where the Gain of GSC is not much smaller than MV, WNMV, and SMV for the deeper sources and (2) the frequency of the DBS signal, where although pulses were at 10 Hz, the pulse-width was 300 μ s and therefore the DBS signal looks very different from the background brain activity and has high frequency harmonics which make it easier for the adaptive algorithm of the GSC (LMS) to remove the background activity and pass DBS pulses. The frequency effect can be seen in Fig. 9, where the Gain of MV, WNMV, and SMV increased by 2 times over the frequency range whereas Gain increased ~ 10 times for GSC.

To substantiate the above explanations, a simulated DBS EEG was created: the DBS component, which was identified by ICA in Section 4.3 was separated from other background activity, then seeded to the same location as that of the DBS electrodes (right thalamus) and superimposed on the background EEG from the same subject. The spatial pass-band was then obtained for the real simulated DBS EEG (11). As can be seen, GSC had the highest Gain of all the beamformers. However, the pass-bands of the MV, WNMV,

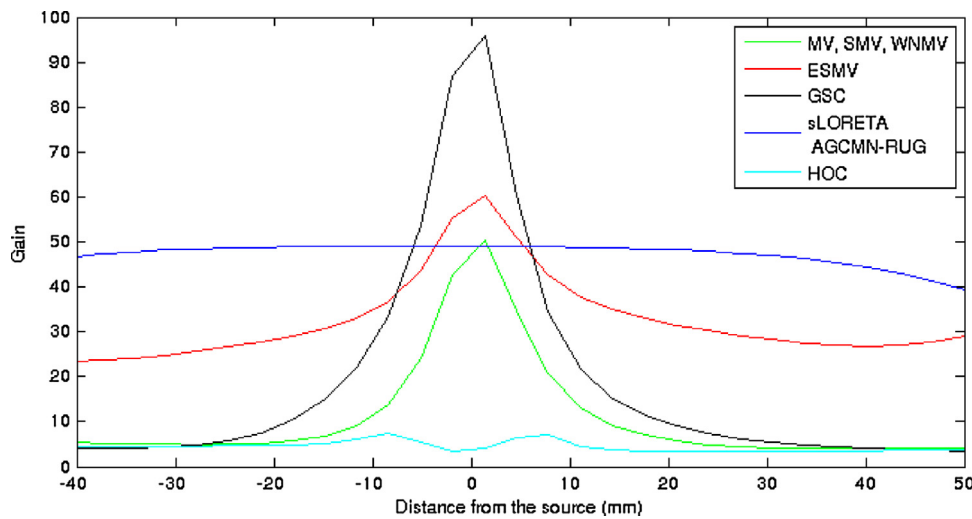


Fig. 11. Spatial pass-bands of the beamformers for the simulated DBS experiment. The simulated DBS source was placed at [8, 14, 1] mm. The time-course of the simulated DBS was the DBS component identified by scalp ICA of the real DBS EEG. To obtain the spatial pass-band, a similar process was applied as described in Fig. 10.

SMV, and GSC for the simulated DBS were much narrower than for the real DBS; i.e., FWHM ~ 14 mm for simulated DBS compared with ~ 40 mm for real DBS. This considerable difference is due to the difference between head modelling (BEM) and the real brain; this is considered one of the main limitations of the application of beamformers in [11]. This difference can also be seen in Fig. 3(d).

5.7. Location bias

The scalar MV, WNMV, SMV, and GSC had uniform white-noise spatial maps. In comparison, sLORETA and AGCMN-RUG had location biases since their white-noise spatial maps were non-uniform (Fig. 12(b)), with the power in some regions being 3 times higher than other regions. This non-uniformity is due to both of these beamformers using the gram matrix $\mathbf{G}_{\Omega}$ which has a matrix norm which varies with the location and, hence, introduces location bias.

The two implementations of the vector beamformers (3-D vector and 3-scalar vector beamformer) resulted in different white-noise spatial maps. However, as the 3-scalar implementation essentially comprises scalar beamformers in each of the 3 orthogonal directions, its white-noise spatial maps are similar to the

scalar version reported above. Hence, for 3-scalar vector implementation, MV, WNMV, SMV, GSC have uniform white-noise spatial map and sLORETA and AGCMN-RUG have location bias with the power 3 times higher in certain regions.

In the case of the 3-D vector beamformers, MV, sLORETA, and AGCMN-RUG have non-uniform white-noise maps, with power in some regions being more than 3 times higher than other regions. This is in line with Eq. 37. GSC also had a non-uniform white-noise spatial map, but to a much smaller extent, with a difference of up to 20%. Only WNMV and SMV had completely uniform white-noise spatial maps. For all the 3-D vector beamformers with non-uniform white-noise spatial maps, the regions with higher power were in locations furthest from the scalp sensors. This location bias for the 3-D vector MV beamformer is similar to Fig. 1(a) and (b) in [10]. As ESMV and HOC are based on MV, their white-noise spatial maps are the same as for MV.

5.8. An important difference between MV, WNMV, and SMV

Although in all simulations, except for location bias, MV, WNMV, and SMV had identical performances, we found an important

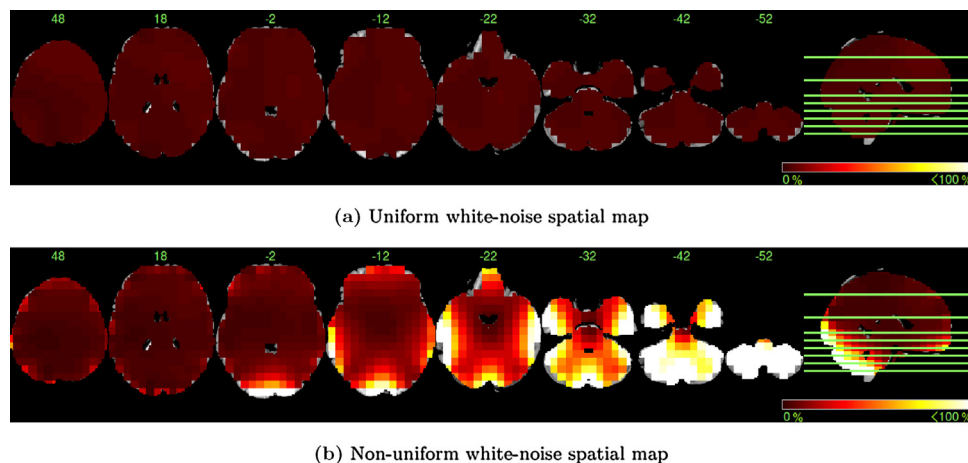


Fig. 12. White-noise spatial maps of the beamformers. The plots show the power differences (%) of the voxels relative to the voxel with the smallest power. (a) A uniform white-noise spatial map in which the power of all voxels are the same, such as for scalar MV, WNMV, SMV, ESMV, HOC, GSC. (b) A non-uniform white-noise spatial map in which some areas have a power more than 3 times to that of other areas, as is the case for sLORETA and AGCMN-RUG. The 3-scalar implementations of the vector beamformers have white-noise spatial maps similar to the scalar beamformers. In contrast, of the 3-D implementations of the vector beamformers, only WNMV and SMV have uniform maps.

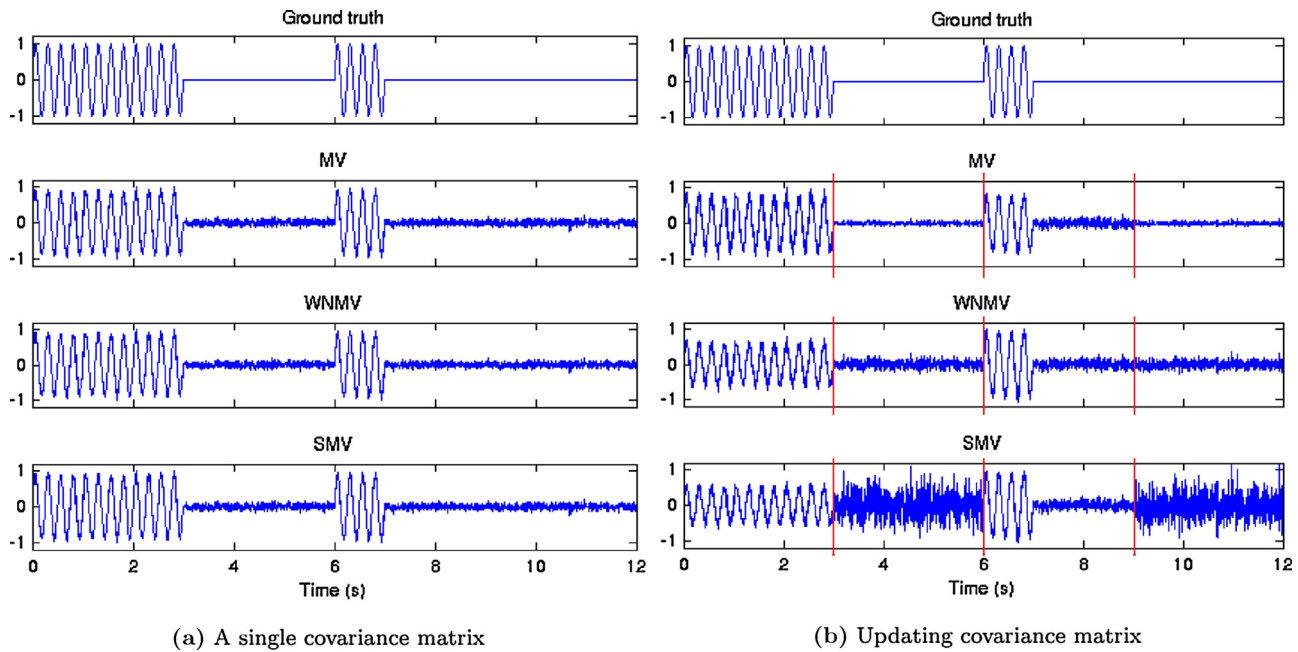


Fig. 13. Time-course reconstruction via MV, WNMV, and SMV for a simulated source (ground truth) in two conditions: (a) a single 12 s covariance matrix, and (b) with a 3 s updating covariance matrix. For ease of demonstration, the background activity was white noise. Updating the covariance matrix introduces visible discontinuity in the reconstructed time-course of the SMV beamformer. The time-courses of SMV and WNMV have magnitude biases as the magnitude of the burst of activity at 6–7 s is different from the preceding burst at 0–3 s for the right plots.

difference between these three beamformers in relation to the covariance matrix. Fig. 13 shows the time-course reconstruction of a simulated source, active from 0–3 s and 6–7 s (ground truth), superimposed on white noise in two situations: the covariance matrix calculated from a 12-s window (the left-hand plots) and the covariance matrix calculated 4 times via a 3-s moving window (the right-hand plots). When the covariance matrix was 12 s, all three beamformers had near identical time-course reconstructions, whereas for the moving 3-s covariance matrix window, SMV magnifies the white noise when there is no source activity. Since in the simulations (Sections 4.1–4.4) the length of the simulated EEG was 12 s (and the simulated source was active at 6–12 s), the covariance matrix has been calculated over the whole 12 s and, therefore, the difference of these three beamformers is minimal. Conversely, if multiple smaller windows are applied, SMV has lower *Gains* in the plots than MV and WNMV. However, other specifications such as spatial and directional pass-bands would not change. Another difference is the magnitude bias: in the left-hand plots, the magnitude of the second burst of activity (6–7 s) is higher than the first burst (0–3 s) for WNMV and SMV, whereas the time-course reconstructed by MV does not have this bias.

6. Discussion

In this study, the performances of 8 beamformers were evaluated and compared based on (1) quality of reconstructed time-courses, (2) spatial and directional pass-bands, (3) effect of simulated source parameters on *Gain* of the beamformers, and (4) white-noise spatial maps. *Gain* was used to evaluate performance and was defined as the ratio of the SNR of the reconstructed time-course to the input SNR. No beamformer was found to be superior in all of the desired characteristics for all conditions. Minimum-variance beamformers had superior spatial and directional resolutions but were more sensitive to changes in source parameters. BEM, which was used for the head modelling in the simulations, cannot take into account the effect on EEG of local changes in skull resistance, which can smear the topography of

sources and especially deep sources. We consider that the results of the simulations are also not only applicable to MEG but are more accurate than EEG due to the smearing influence of the skull being minimal for MEG.

The behaviour of each beamformer is summarized and discussed below.

6.1. Minimum-variance beamformers

The minimum-variance (MV), weight-normalized minimum-variance (WNMV), standardized minimum-variance (SMV), eigen-space minimum-variance (ESMV), and higher-order covariance matrix (HOC) beamformers all belong to the minimum-variance class of beamformers. MV, WNMV, and SMV had identical performances in terms of sensitivity to changes in the input parameters and spatial and directional pass-bands. ESMV performed similarly to the other minimum-variance beamformers except for having wider spatial and directional pass-bands and a higher *Gain* due to the removal of noise space from the covariance matrix. However, this advantage is only present when sources of interest are strong. A similar limitation of subspace-projection-based spatial filters was reported in [16]. Therefore, ESMV is most suitable for strong source signals such as large evoked potentials [43]. Similarly, ESMV performs better in reconstruction of shallower source signals, as they are spatially more separable than deep EEG sources [35].

With respect to changes in input parameters, the minimum-variance beamformers (MV, WNMV, SMV, ESMV) were shown to be much more sensitive to changes in depth, magnitude, and the frequency of the simulated sources than sLORETA and AGCMN-RUG. Conversely, the minimum-variance beamformers are less sensitive to changes in source orientation than sLORETA and AGCMN-RUG. Through the simulations and the DBS study, the spatial pass-band of MV, WNMV, and SMV were found to be superior to sLORETA and AGCMN-RUG but not better than GSC. The directional pass-band of MV, WNMV, and SMV were superior to the other beamformers. In the case of ESMV, the DBS experiment showed that their spatial and directional pass-bands are not as good as the other

minimum-variance beamformers (MV, WNMV, and SMV). MV, WNMV, SMV, and ESMV maintained higher *Gain* than the other beamformers in all situations except the DBS experiment, due to the narrow pulses of the DBS which have high-frequency harmonics which enables GSC to achieve a higher *Gain*.

Although MV, WNMV, and SMV performed identically with changes in input parameters and identical pass-bands, there are differences between these beamformers. In the case of white-noise spatial map, the scalar versions of MV, WNMV, and SMV and, hence, their 3-scalar vector implementation, had uniform white-noise spatial maps when the normalized lead-field was used. But for the 3-D implementation of the vector version, MV had a non-uniform white-noise spatial map and, therefore, had a location bias, whereas WNMV and SMV had uniform white-noise spatial maps. Another difference between MV, WNMV, and SMV is magnitude bias, where updating the covariance matrix introduces non-uniformity in the magnitude of the reconstructed time-courses of SMV in particular.

The HOC beamformer had very poor time-course reconstruction under all conditions. In fact, there was no example in the original paper [9] of time-series reconstruction by this beamformer, as it was only compared to other beamformers in terms of neural activity index. We included the HOC in the current study as it had been compared with other beamformers, including MV and WNMV, and shown to be superior for source localization by neural activity index for strong sources [9]. However, HOC achieves a higher SNR than MV via the neural activity index at each time sample, which was the focus in [9], whereas our focus was on performance of beamformers on the reconstruction of the actual time-courses of sources.

6.2. Minimum-norm beamformers

In this study, three minimum-norm beamformers were evaluated: standardized low-resolution electromagnetic tomography (sLORETA), array-gain constraint minimum-norm filter with recursively updating gram matrix (AGCMN-RUG), and the generalized sidelobe canceller form of the quiescent beamformer (GSC). Throughout this study, sLORETA had an identical performance to AGCMN-RUG. sLORETA and AGCMN-RUG were found to be linear with respect to magnitude and frequency of the simulated sources, but had changes in *Gain* with respect to changes in the orientation and depth of the source. They had average *Gains* several times smaller than the minimum-variance beamformers for small weak sources, but similar *Gains* for strong sources. sLORETA and AGCMN-RUG were also found to have poor spatial and directional resolutions compared with minimum-variance beamformers. Another important point with these two beamformer is that they had errors identifying actual orientations of the sources in both the simulation and the DBS experiment. Both of these beamformers had wider spatial pass-bands than directional pass-bands. AGCMN-RUG was proposed in [13] and it was stated AGCMN-RUG does not need a covariance matrix of the sensor data. However, in this context the covariance matrix was applied for regularization of the gram matrix. The original paper stated that AGCMN-RUG cannot replace the minimum-variance beamformer as it has a spatial resolution intermediate to sLORETA and MV.

GSC performed similar to the minimum-variance beamformers (e.g., nonlinear to magnitude and frequency of the simulated source) and had a narrower spatial pass-band than sLORETA. However, the directional pass-band of GSC was widest both in the simulation and in the DBS experiment. Overall, GSC had lower *Gains* to those of the minimum-variance beamformers. However, an important exception to this was seen in the DBS experiment, in which the short pulse width of the DBS signals contained high-frequency harmonics which made it easier for the adaptive LMS algorithm of the GSC to filter the DBS pulses from the background activity.

7. Conclusion

The aim of the current study was to examine and compare the performance of 8 beamformers from the minimum-variance and minimum-norm families of spatial filters, particularly with respect to variations in depth, orientation, magnitude, and frequency of the simulated sources. Spatial and directional pass-bands and white-noise spatial maps of these beamformers were also investigated. A real source signal from a DBS patient was also investigated for comparison with results from simulated sources. *Gain* was used to describe the behaviour of the beamformers and was defined as the ratio of the SNR of the reconstructed time-course to that of the input signal on EEG. In summary, minimum-variance beamformers have higher *Gains*, although their performance is more affected by changes in magnitude, depth, and frequency of source signals. Minimum-norm beamformers have lower *Gains* but are more invariant to changes in magnitude, depth, and frequency of the sources. Minimum-variance beamformers have superior spatial resolution. The white-noise spatial maps of the beamformers show that some vector beamformers (WNMV, SMV) have no location bias irrespective of whether implemented as 3-D vector or 3-scalar vector beamformers, whereas others always have bias (sLORETA, AGCMN-RUG) and some (MV and GSC) only have bias for 3-scalar vector implementation.

Acknowledgements

This work was supported by University of Otago by way of a Postgraduate Scholarship and by Marsden Fund by way of a Neurotechnology Scholarship. The funding sources had no role in the design or conduction of the study, collection, management, analysis, or interpretation of the data.

The authors thank Simon Knopp for his suggestions and thoughts on this work.

References

- [1] J.H. Hong, S.C. Jun, Scanning reduction strategy in MEG/EEG beamformer source imaging, *J. Appl. Math.* 2012 (2012) 1–19.
- [2] M. Besserve, K. Jerbi, F. Laurent, S. Baillet, J. Martinerie, L. Garnero, Classification methods for ongoing EEG and MEG signals, *Biol. Res.* 40 (2007) 415–437.
- [3] P. Poolman, R.M. Frank, P. Luu, S.M. Pederson, D.M. Tucker, A single-trial analytic framework for EEG analysis and its application to target detection and classification, *NeuroImage* 42 (2008) 787–798.
- [4] M. Congedo, F. Lotte, A. Lecuyer, Classification of movement intention by spatially filtered electromagnetic inverse solutions, *Phys. Med. Biol.* 51 (2006) 1971–1989.
- [5] B. Blankertz, R. Tomioka, S. Lemm, M. Kawanabe, K. Muller, Optimizing spatial filters for robust EEG single-trial analysis, *IEEE Signal Process. Mag.* 25 (2008) 41–56.
- [6] M. Ahn, J.H. Hong, S.C. Jun, Source space based brain computer interface, in: *Proc. 17th IFMBE Int. Conf. Biomed.*, 2010, pp. 366–369.
- [7] K. Sekihara, S.S. Nagarajan, D. Poeppel, A. Marantz, Y. Miyashita, Reconstructing spatio-temporal activities of neural sources using an MEG vector beamformer technique, *IEEE Trans. Biomed. Eng.* 48 (2001) 760–771.
- [8] B.D. Van Veen, K.M. Buckley, Beamforming: a versatile approach to spatial filtering, *IEEE Mag. Acoust. Speech Signal Process.* 5 (1988) 4–24.
- [9] M.X. Huang, J.J. Shih, R.R. Lee, D.L. Harrington, R.J. Thoma, M.P. Weisend, F. Hanlon, K.M. Paulson, T. Li, K. Martin, G.A. Millers, J.M. Canive, Commonalities and differences among vectorized beamformers in electromagnetic source imaging, *Brain Topogr.* 16 (2004) 139–158.
- [10] B.D. Van Veen, W. van Drongelen, M. Yuchtman, A. Suzuki, Localization of brain electrical activity via linearly constrained minimum variance spatial filtering, *IEEE Trans. Biomed. Eng.* 44 (1997) 867–880.
- [11] M.J. Brookes, K. Mullinger, C. Stevenson, P. Morris, R. Bowtell, Simultaneous EEG source localisation and artifact rejection during concurrent fMRI by means of spatial filtering, *NeuroImage* 40 (2008) 1090–1104.
- [12] R.E. Greenblatt, A. Ossadtchi, M.E. Pflieger, Local linear estimators for the linear bioelectromagnetic inverse problem, *IEEE Trans. Biomed. Eng.* 53 (2005) 3403–3412.
- [13] I. Kumihashi, K. Sekihara, Array-gain constraint minimum-norm spatial filter with recursively updated gram matrix for biomagnetic source imaging, *IEEE Trans. Biomed. Eng.* 57 (2010) 1358–1365.

- [14] D.M. Ward, R.D. Jones, P.J. Bones, G.J. Carroll, Enhancement of deep epileptiform activity in the EEG via 3-D adaptive spatial filtering, *IEEE Trans. Biomed. Eng.* 46 (1999) 707–716.
- [15] J.P. Owen, D.P. Wipf, H.T. Attias, K. Sekihara, N.S.S., Performance evaluation of the Champagne source reconstruction algorithm on simulated and real M/EEG data, *Neuroscience* 60 (2012) 305–323.
- [16] M. Congedo, Subspace projection filters for real-time brain electromagnetic imaging, *IEEE Trans. Biomed. Eng.* 53 (2006) 1624–1634.
- [17] M. Grosse-Wentrup, C. Liefhold, K. Gramann, M. Buss, Beamforming in noninvasive brain-computer interfaces, *IEEE Trans. Biomed. Eng.* 56 (2009) 1209–1219.
- [18] S.E. Robinson, J. Vrba, Functional neuroimaging by synthetic aperture magnetometry (SAM), in: T. Yoshimoto, M. Kotani, S. Kuriki, H. Karibe, N. Nakasato (Eds.), *Proc. 11th Int. Conf. Biomagnetism*, Tohoku Univ. Press, Sendai, 1998, pp. 302–305.
- [19] J. Vrba, S.E. Robinson, Signal processing in magnetoencephalography, *Methods* 25 (2001) 249–271.
- [20] S. Johnson, G. Prendergast, M. Hymers, G. Green, Examining the effects of one- and three-dimensional spatial filtering analyses in magnetoencephalography, *PLoS ONE* 6 (2011) e22251.
- [21] Y. Jonmohamadi, G. Poudel, C. Innes, R. Jones, Voxel-ICA for reconstruction of source signal time-series and orientation in EEG and MEG, *Aust. Phys. Eng. Sci. Med.* 37 (2014) 457–464.
- [22] G.V. Borgia, L.J. Kaplan, Superresolution of uncorrelated interference sources by using adaptive array techniques, *IEEE Trans. Antennas Propag.* AP-27 (6) (1979) 842–845.
- [23] J.L. Yu, C.C. Yeh, Generalized eigenspace-based beamformers, *IEEE Trans. Signal Process.* 43 (1995) 2453–2461.
- [24] M. Spencer, R.M. Leahy, J.C. Mosher, P.S. Lewis, Adaptive filters for monitoring localized brain activity from surface potential time series, *Proc. IEEE Asilomar Conf. Signal Syst. Comput.* 26 (1992) 156–160.
- [25] R.D. Pascual-Marqui, Standardized low-resolution brain electromagnetic tomography (sLORETA): technical details, *Methods Find. Exp. Clin. Pharmacol.* 24 (Suppl. D) (2002) 5–12.
- [26] M.J. Brookes, J. Vrba, S. Robinson, C.M. Stevenson, A. Peters, G. Barnes, A. Hillebrand, P.G. Morris, Optimising experimental design for MEG beamformer imaging, *NeuroImage* 39 (2008) 1788–1802.
- [27] J. Gross, A.A. Ioannides, Linear transformations of data space in MEG, *Phys. Med. Biol.* 44 (1999) 2081–2097.
- [28] H. Cox, R. Zeskind, M. Owen, Robust adaptive beamforming, *IEEE Trans. Acoust. Speech Signal Process.* ASSP-35 (1987) 1365–1376.
- [29] A. Hillebrand, G.R. Barnes, Beamformer analysis of MEG data, *Int. Rev. Neurobiol.* 68 (2005) 149–171.
- [30] C.P. Hansen, Rank-Deficient and Discrete Ill-Posed Problems: Numerical Aspects of Linear Inversion, SIAM, Philadelphia, 1998.
- [31] K. Sekihara, M. Sahani, S.S. Nagarajan, Localization bias and spatial resolution of adaptive and non-adaptive spatial filters for MEG source reconstruction, *NeuroImage* 25 (2005) 1056–1067.
- [32] K. Sekihara, S.S. Nagarajan, Adaptive Spatial Filters for Electromagnetic Brain Imaging, Springer, Berlin, 2008.
- [33] H. Dang, K. Ng, J. Kroger, Novel beamformers for multiple correlated brain source localization and reconstruction, in: *IEEE Int. Conf. Acoust. Speech. Sig. Process.*, 2011, pp. 721–724.
- [34] R.O. Schmidt, A Signal Subspace Approach to Multiple Emitter Location and Spectral Estimation, Ph.D. thesis, Stanford University, 1981.
- [35] K. Sekihara, S.S. Nagarajan, D. Poeppel, A. Marantz, Y. Miyashita, Application of an MEG eigenspace beamformer to reconstructing spatio-temporal activities of neural sources, *Hum. Brain Mapp.* 15 (2002) 199–215.
- [36] T.F. Oostendorp, A. van Oosterom, Source parameter estimation in inhomogeneous volume conductors of arbitrary shape, *IEEE Trans. Biomed. Eng.* 36 (1989) 382–391.
- [37] R. Oostenveld, P. Fries, E. Maris, J.M. Schoffelen, Fieldtrip: open source software for advanced analysis of MEG, EEG, and invasive electrophysiological data, *Comput. Intell. Neurosci.* 2011 (2011) 1–9.
- [38] A.J. Bell, T.J. Sejnowski, An information-maximization approach to blind separation and blind deconvolution, *Neural Comput.* 7 (1995) 1129–1159.
- [39] A. Delorme, S. Makeig, EEGLAB: an open source toolbox for analysis of single-trial EEG dynamics, *J. Neurosci. Meth.* 134 (2004) 9–21.
- [40] S. Makeig, S. Debener, J. Onton, A. Delorme, Mining event-related brain dynamics, *Trends Cogn. Sci.* 8 (2004) 204–210.
- [41] Y. Jonmohamadi, G. Poudel, C. Innes, R. Jones, Electromagnetic tomography via source-space ICA, *Proc. Int. Conf. IEEE Eng. Med. Biol. Soc.* 35 (2013) 37–40.
- [42] Y. Jonmohamadi, G. Poudel, C. Innes, R. Jones, Source-space ICA for EEG source separation, localization, and time-course reconstruction, *NeuroImage* (2014), <http://dx.doi.org/10.1016/j.neuroimage.2014.07.052> (in press).
- [43] S.S. Dalal, J.M. Zumer, A.G. Guggisberg, M. Trumpis, D.D. Wong, K. Sekihara, S.S. Nagarajan, MEG/EEG source reconstruction, statistical evaluation, and visualization with NUTMEG, *IEEE Trans. Biomed. Eng.* 2011 (2011) 1–17.