

Automatic ocular artifacts removal in EEG using deep learning

Banghua Yang*, Kaiwen Duan, Chengcheng Fan, Chenxiao Hu, Jinlong Wang

Department of Automation, School of Mechatronics Engineering and Automation, Key Laboratory of Power Station Automation Technology, Shanghai University, Shanghai 200072, China



ARTICLE INFO

Article history:

Received 13 April 2017

Received in revised form 12 January 2018

Accepted 24 February 2018

Available online 23 March 2018

Keywords:

Ocular artifacts (OAs) removal
Electroencephalogram (EEG)
Deep learning network (DLN)
Independent component analysis (ICA)
Shallow network

ABSTRACT

Ocular artifacts (OAs) are one the most important form of interferences in the analysis of electroencephalogram (EEG) research. OAs removal/reduction is a key analysis before the processing of EEG signals. For classic OAs removal methods, either an additional electrooculogram (EOG) recording or multi-channel EEG is required. To address these limitations of existing methods, this paper investigates the use of deep learning network (DLN) to remove OAs in EEG signals. The proposed method consists of offline stage and online stage. In the offline stage, training samples without OAs are intercepted and used to train an DLN to reconstruct the EEG signals. The high-order statistical moments information of EEG is therefore learned. In the online stage, the trained DLN is used as a filter to automatically remove OAs from the contaminated EEG signals. Compared with the exiting methods, the proposed method has the following advantages: (i) nonuse of additional EOG reference signals, (ii) any few number of EEG channels can be analyzed, (iii) time saving, and (iv) the strong generalization ability, etc. In this paper, both public database and lab individual data for EEG analysis are used, we compared the proposed method with the classic independent component analysis (ICA), kurtosis-ICA (K-ICA), Second-order blind identification (SOBI) and a shallow network method. Experimental results show that the proposed method performs better even for very noisy EEG.

© 2018 Elsevier Ltd. All rights reserved.

1. Introduction

Electroencephalography (EEG) is a portable and high-temporal resolution signal that can be used for quantitative analysis of the brain's different functional states [1–7]. Typically, EEG signals are characterized by three components, including shape, frequency, and amplitude [8]. In the frequency domain, EEG signals usually implicate information about rhythmic activities at different frequency bandwidths of δ -delta (0.5–4 Hz), θ -theta (4–8 Hz), α -alpha (8–13 Hz), β -beta (13–30 Hz) and γ -gamma (30–50 Hz) [8–11]. Main applications of EEG range from neurophysiological clinical monitor to brain computer interface (BCI) [12–17]. However, a recorded EEG signal is highly contaminated with physiological artifacts from various sources, such as eye blinking/movements, heart beating and movement of other muscle groups [8,18]. Among the artifacts signals, one of the most troublesome artifacts observed in the EEG signals are due to the ocular activity, which are called ocular artifacts (OAs) [19]. OAs can easily obscure the EEG signal, making its visual or automated neurophysiological monitor and data analy-

sis difficult [20]. Moreover, it has yet been proven that OAs diminish the classification accuracy of BCI applications [1,21].

EEG signals contain lots of OAs. During the process of EEG signals acquisition, eye blinking/movements are inevitable. Eye blinking/movements generate spike-like shaped signal waveforms with the peaks reaching up to 800 μ V and each peak occurs in a very short period of 200–400 ms [22], which are called electrooculogram (EOG) signals. These movements change the electric field distribution around the eyes and the surface of brain cortex. When these movements are collected by the scalp EEG electrode, the OAs are formed and overlapped with EEG signals. EOG signals share many features with OAs such as the low frequency bands (δ , θ and α [23]), spike-like shape, etc. and they are all occurring in a very short period. Estimating clean EEG signals from contaminated ones is a very tricky problem, because OAs are unmeasurable interferences. Fortunately, the amplitudes of OAs are larger than that of EEG, which may be a good breakthrough for OAs removal.

OAs removal is inevitable before the neurophysiological monitoring or data analysis of EEG. It's true that the all kinds of algorithms aiming at recovering artifact-free signal have been investigated intensively. In the early years, the common method was to manually cut the entire segment of data affected by the OAs, which can lead to a relevant information loss [1,13]. In recent

* Corresponding author.

E-mail address: yangbanghua@shu.edu.cn (B. Yang).

years, there has been an increasing interest in the research of OAs removal [24–27]. It is the evidence to the importance of this issue. Generally, these methods can be divided into four categories.

- (1) **Artifacts filters:** This class of methods are common artifacts removal methods for EEG signals, which have been generally used to remove OAs. One of the most popular methods is the linear filtering method [12,28]. This method assumes that the OAs and EEG signals have different frequency ranges. When removing OAs, a high-pass filter can be used to filter the OAs in EEG signals. This method can be efficient for BCI systems which use a neurological phenomenon with high-frequency bands (like β and γ rhythms). However, some relevant research works have proven that the OAs are overlapped with the clean EEG signals in frequency domain [8,29]. Therefore, it might eliminate useful EEG signals during artifacts removal. The adaptive filtering method can be seen as an improved method of the linear filtering method. This method employs the contaminated EEG signal as input, and the EOG signals as reference signals. The corrected EEG can be obtained by subtracting the filtered outputs [30,31]. This method is more effective compared with the linear filtering method. However, it needs an additional EOG recording. Unfortunately, collecting EOG signals during a long-term EEG recording is inconvenient and uncomfortable for the subjects. As the modern brain computer interface rehabilitation system develops rapidly, the producers are considering the good user's experience as one of the top priorities. For the purpose of putting motor imagery system into practical daily rehabilitation for stroke patients, some patients could not accept the electrode "paste" on their head, not even the EOG electrode which is around their eyes. So, it's better to use few EEG channels and not to use EOG channel to denoise the OAs to a satisfying result, which is very important to practical using of brain computer interfaces. One of our research goal is helping patients acceptable to use brain computer interfaces with good user experience.
- (2) **Blind source separation (BBS):** This class of methods have been shown by many researchers that they can be used to efficiently separate the distinct artifact components from EEG data [36]. Blind source separation assumes that the EEG signals are a linear mixture of neural and OAs. Independent component analysis (ICA) is one of most popular blind source separation methods [37,38]. ICA can decompose multiple EEG channels into an equal number of independent components (ICs). ICs that contain OAs can be identified through visual inspection. After discarding OA-related ICs, the uncorrupted EEG signals can be reconstructed by the remaining ICs. ICA can separate OA-related ICs efficiently, but this method not only causes useful information loss but also enhances the coherence between different EEG channels [39]. In addition, this method needs a large number of EEG channels to ensure that the ICs with distinct characteristics of OAs are identified. However, collecting large number of EEG channels is also uncomfortable for the subject. With few number of EEG analyzing channels and low time cost, combining without additional EOG channel recording, the method proposed in the paper could be used in practical motor imagery rehabilitation system, which avoid the trivial preparation and experiment and more susceptible for stroke patients.
- (3) **Wavelet transformation analysis:** This class of methods have been widely used in OAs removal [32,33]. In particular, the wavelet analysis that based on thresholding techniques has received extensive attentions [34,35]. This method uses wavelet transformation to decompose the contaminated EEG signals into series of sub-bands at first. Then a thresholding function is used to automatically correct the coefficients related to OAs sub-bands. Finally, EEG signals are reconstructed based

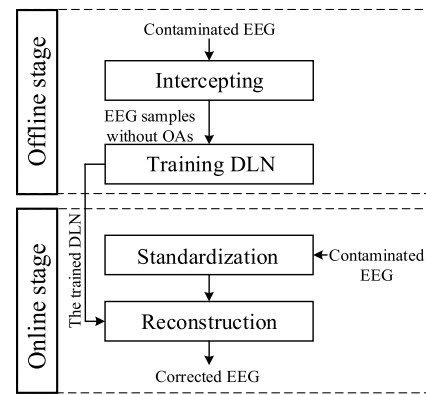


Fig. 1. Flowchart of the proposed method.

on the corrected coefficients. This method can be used in the online OAs removal and it doesn't need reference signals. However, the threshold for wavelet transformation is difficult to be determined, an unsuitable threshold may result in degradation of the EEG data which could reduce the generalization ability of the EEG system [26].

- (4) **Neural network analysis:** Recently, the neural network has been applied in the OAs artifacts removal [8]. Nguyen H.A.T et al. proposed a method called wavelet neural network to remove the OAs artifacts [8]. This method replaced the thresholding function that used in the wavelet analysis with a neural network, which avoids the defects of thresholding function. Jing Hu et al. combined the functional link neural network and adaptive neural fuzzy inference system (FLNN-ANFIS) to remove OAs and electromyogram (EMG) artifacts [26]. This method wisely constructed a filter by using FLNN and ANFIS to filter OAs and EMG artifacts, which is good at handing the vague data. These two methods mentioned above both take advantage of the ability of neural network that can approximate smooth nonlinear functions. However, these methods still need the EOG signals to train the neural network or these methods still need EOG signals as reference signals. Moreover, the approximation ability of neural network is limited compared with the deep learning network.

As we all know, deep learning network (DLN) has been widely used in image processing and other applications. Despite the success of DLN, its application in EEG is still rare [40]. In an attempt to overcome defects of the traditional methods mentioned above, this paper proposes a novel, robust and efficient DLN to remove OAs in contaminated EEG. The proposed method has the following properties: (i) this method doesn't need additional EOG recording as reference signals whether in offline stage or online stage, which is comfortable for the subjects with foreseeable good application on brain computer interface rehabilitation system, (ii) this method is suitable for few number of EEG electrodes, which is convenient for EEG recording, cost-saving and appropriate for application, (iii) OAs can be automatically removed online by using the learned model, which has fast processing speed and it can be applied in online system, (iv) this method has strong generalization ability, which can be widely used. The proposed method can be divided into two stages (Fig. 1). The first stage is an offline stage, and the second stage is an online stage. A formal flowchart description is given in Table 1.

The rest of the paper consists of four sections. Section 2 introduces the proposed method and related work. Section 3 gives the description of the source of the dataset used in this study. Section 4 shows the detailed experiments and results. Conclusion and discussion are given in Section 5.

Table 1
Component of the proposed method.

Offline stage:

- (i): provide contaminated EEG and intercept enough training samples.
- (ii): use these samples to train an DLN to learn features of clean EEG.

Online stage:

- (iii): standardize the contaminated EEG and input them to the trained DLN.
- (iv): apply the trained DLN to reconstruct the input data and output the corrected EEG.

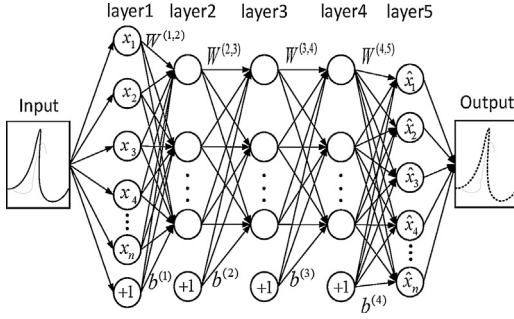


Fig. 2. Structure of DLN.

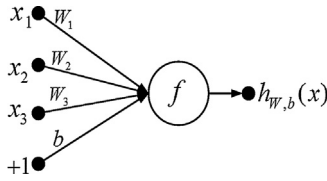


Fig. 3. A single "neuron".

2. Materials and methods

2.1. Proposed method

2.1.1. Introduction of DLN

DLN is the new family of learning methods that can offer good representation of data using a multiple-layered structure, with each layer representing different degree of data features abstraction [41–44]. In this study, we apply the stacked sparse autoencoder (SSAE) [45] as a DLN. Fig. 2 shows a DLN with 3 hidden layers. Assuming that $x, \hat{x}$ are input data and output data of the DLN. According to [45], it uses $x = \hat{x}$, where $x, \hat{x}$ are n dimensional column vectors, which correspond to n inputs of DLN. Layer 1 is the input layer, layer 2–4 are hidden layers, layer 5 is the output layer. The circles labeled "+1" are called bias units and correspond to the intercept terms, while other circles are called "neurons". Fig. 3 is a single "neuron". In Fig. 3, W_1-W_3 is the connection weights and b is the intercept term. The "neuron" is an arithmetic unit with x_1-x_3 and b as the input. The output can be written as follows:

$$h_{w,b}(x) = f(W^T x) = f\left(\sum_{i=1}^3 W_i x_i + b\right) \quad (1)$$

where $h_{w,b}(x)$ is called the activation value and $f(\cdot)$ is called the activation function. In this paper, we use the sigmoid function as the activation function, which is:

$$f(z) = \frac{1}{1 + \exp(-z)} \quad (2)$$

where $f(z) \in (0, 1)$.

As shown in Fig. 2, $W^{(l,l+1)}$ ($l = 1, 2, 3, 4$) is the parameter (or weight) matrix associated with layer l and layer $l+1$. Also, $b^{(l)}$

denotes the bias associated with layer l and layer $l+1$. Two adjacent layers are connected by weights and intercept terms. The structure of the DLN can be seen as a deep neural network consisting of multiple layers of sparse autoencoders (SAE), in which the output of former SAE is wired to the input of successive SAE. Following the method of [45], a good way to obtain good parameters for a DLN is to use greedy layer-wise training [46]. To do this, (i) train the first SAE on original input to obtain parameters accordingly, (ii) use the first layer of the first SAE to transform the original input into a vector consisting of activation value of the hidden units, (iii) cut off the output layer and the connection weights associated with hidden layer and output layer in the SAE, (iv) train the second SAE on this vector to obtain parameters of the second SAE accordingly, (v) repeat for subsequent layers, using the output of each layer as input for the subsequent layer. The training process can be seen in Fig. 4. The red part in Fig. 4 means that it will be eliminated in the next SAE training. After this phase of training is complete, fine-tuning using backpropagation can be used to improve the results by tuning the parameters of all layers at the same time. The approach about how to train an SAE can be seen in [47].

2.1.2. Interception of EEG samples

In this paper, we use the EEG samples without OAs to train an DLN. By doing this step, we could successfully avoid the use of EOG signals, which is one of the advantages of the proposed method. So, the way to intercept the ideal training samples is the key step of this paper. For realizing this, it's necessary to obtain the EEG segments without OAs at first. For convenience's sake, the EEG segments without OAs will be simplified as no OAs EEG segments (NOEs) and the EEG segment without OAs as NOE.

(1) Introduction of kurtosis

Kurtosis is a fourth-order cumulant of a random variable. Kurtosis is defined by

$$\text{Kur} = cm_4 - 3cm_2^2 \quad (3)$$

where cm_j ($j = 2, 4$) is the j^{th} order central moment with the following form:

$$cm_j = E\{[s - E(s)]^j\} \quad (4)$$

where E is the statistical expectation of the random variable s . In a signal, the steeper a signal is, the higher kurtosis value of the signal will be. Thus, the kurtosis is a good method to detect OAs.

(2) Identification of NOEs based on kurtosis

As is mentioned in the introduction part, EOG signals share many features with OAs, for example, they both occur in a very short period (200–400 ms). In other words, there exists NOEs in a contaminated EEG signal. Fig. 5 is a typical contaminated EEG signal for 20 s (the sampling frequency was 100 Hz). The signal is divided with the same time period into 10 windows. Fig. 5 shows that OAs mainly exist in the Window 3, Window 4, Window 8 and Window 9, while the EEG segments in the rest windows can be considered as not containing OAs. Calculate the kurtosis of the EEG segment in each window and obtain the kurtosis value k_w (w denotes the sequence number of window). Then according to a certain amount of observation and empirical analysis, threshold k_v is set to pick up the EEG segment. If the EEG segment kurtosis is over threshold k_v , this segment is eliminated which is doom to contain ocular artifact. On the contrary, the rest remaining EEG segments are considered as non-contaminated EEG segments. In this way, we can obtain enough NOEs.

(3) Interception of the training samples from NOEs

In this paper, we intercept the training samples from NOEs. Assume that the length of a window is L_w , which is also the length of an NOE. Assume that the length of a training sample is n . α is the parameter for letting the L_w and n to be integral multiple relationship. Due to the α which is the positive real number parameter, all

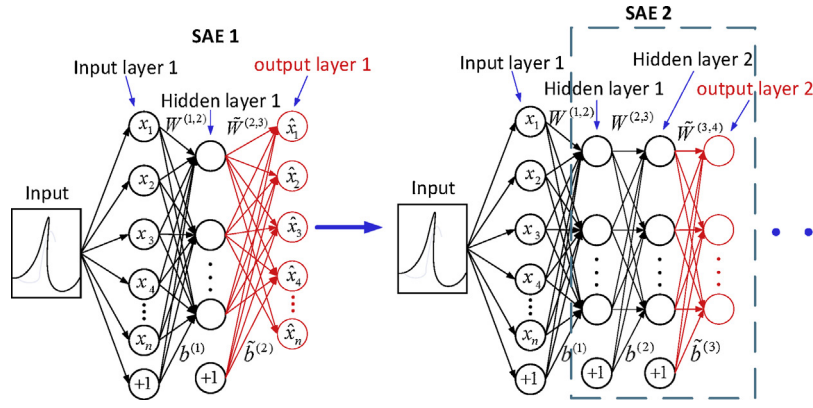


Fig. 4. Schematic of training process.

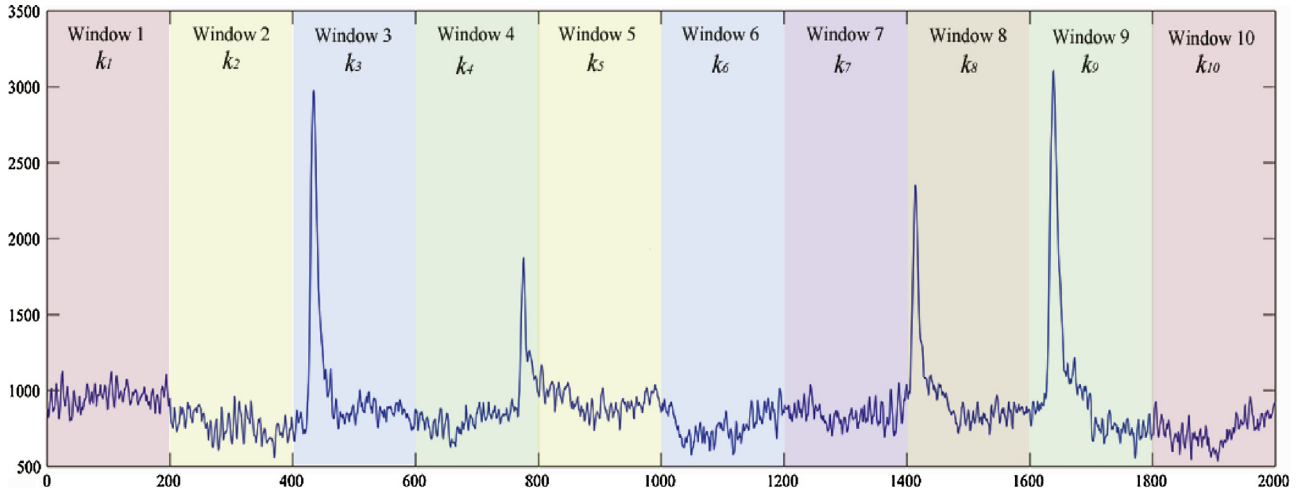


Fig. 5. A typical contaminated EEG signal for 20s.

the length of the intercepted segments is equal. The alignment of the intercepted segment could make full use of the EEG segment and it has benefit to the further analysis. In order to make full use of the selected NOEs, the relationship of L_w and n should satisfy:

$$L_w = \alpha n (\alpha \in \mathbb{N}_+) \quad (5)$$

2.1.3. EEG feature learning by DLN training during offline stage

The DLN needs to be trained firstly before employing, which is one of the most important content of this method. A training set of m training samples $\{x^{(1)}, \dots, x^{(m)}\}$ is firstly intercepted from NOEs. Where every sample $(x^{(a)} (a \in \mathbb{N}_+, a \leq m))$ is n dimensional column vector (contains n sampling points), which corresponds to n inputs of DLN. Each training sample should be normalized before training. A formal description of the framework for 3 hidden layers is given in Table 2.

After the DLN training described in Table 2, the learnt optimal parameter set θ is computed. That is also to say that some features of EEG signals have been learned by these parameters. The next stage is an online stage, the contaminated EEG signal should be standardized before OAs removal.

2.1.4. Standardization of contaminated EEG before OAs removal

The contaminated EEG should be standardized before OAs removal, which is a very crucial step. Select an NOE from the contaminated EEG and normalize it at first. Then take this segment as a reference and standardize the rest of contaminated EEG according to the reference. Assume that the length of a contaminated EEG seg-

ment is L_c (contains L_c sampling points), then the standardization formula is as follows:

$$EEG_{std}(k) = \frac{(EEG_{ctn}(k) - SN_{min})}{SN_{max} - SN_{min}} - \min \left\{ \frac{(EEG_{ctn}(x) - SN_{min})}{SN_{max} - SN_{min}} \mid x \in (1, 2, \dots, L_c) \right\} \quad (6)$$

where $EEG_{std}(k)$ denotes the contaminated EEG after standardization on the sampling point k , EEG_{ctn} denotes the contaminated EEG before standardization, SN_{min} and SN_{max} denote the minimum and maximum of the selected NOE, respectively. The anti-standardization formula is as follows:

$$EEG_{crt}(k) = \left[EEG_{otp}(k) + \min \left\{ \frac{(EEG_{ctn}(x) - SN_{min})}{SN_{max} - SN_{min}} \mid x \in (1, 2, \dots, L_c) \right\} \right] \cdot (SN_{max} - SN_{min}) + SN_{min} \quad (7)$$

where EEG_{crt} denotes the corrected EEG, EEG_{otp} denotes the output of the DLN. Fig. 6 and 7 show a reference for 1s and an EEG_{std} for 2s, respectively. From Fig. 7 we can see that the NOEs in the EEG_{std} are successfully restricted to 0–1, while the OAs are greater than 1. Since we use Eq. (2) as the activation function ($f(z) \in (0, 1)$), the output range of a “neuron” is 0–1, which means the output range of the output layer composed of “neurons” in a DLN will also be 0–1. And since we use the normalized EEG samples without OAs to train the DLN, the NOEs will be reconstructed and the OAs will be limited when the trained DLN is applied to remove OAs in EEG_{std} . In this way, the OAs in the contaminated EEG can be successfully removed,

Table 2
DLN training algorithm.

Input: A fixed training set of m training samples $\{x^{(1)}, \dots, x^{(m)}\}$.

Output: The learnt optimal weight and bias θ , $\theta = (W^{(1,2)}, b^{(1)}, W^{(2,3)}, b^{(2)}, W^{(3,4)}, b^{(3)}, W^{(4,5)}, b^{(4)})$.

```

1: for  $i = 1$  to 3 do
2:   Train the  $i^{th}$  SAE model
3:   if  $i < 3$ 
4:     Obtain optimal parameters  $\hat{W}^{(i,i+1)}, \hat{b}^{(i)}$ .
5:     Obtain a vector consisting of activation of the hidden units.
6:     Use this vector as input for the subsequent SAE.
7:   else
8:     Obtain optimal parameters  $\hat{W}^{(3,4)}, \hat{b}^{(3)}, \hat{W}^{(4,5)}, \hat{b}^{(4)}$ .
9: end for

10: Learn an optimal parameter set  $\hat{\theta} = (\hat{W}^{(1,2)}, \hat{b}^{(1)}, \hat{W}^{(2,3)}, \hat{b}^{(2)}, \hat{W}^{(3,4)}, \hat{b}^{(3)}, \hat{W}^{(4,5)}, \hat{b}^{(4)})$ .

11: Obtain the optimal parameter set  $\theta$  by fine-tuning.

12: return the learnt optimal parameter set  $\theta$ .
```

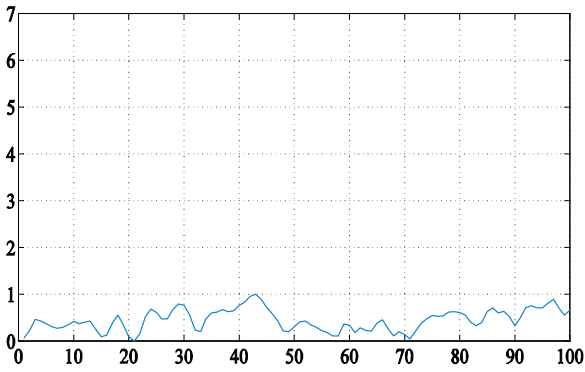


Fig. 6. A reference for 1s.

and this is the core of the proposed method. After the standardization of contaminated EEG, the trained DLN can be applied to remove OAs.

2.1.5. OAs removal by using the trained DLN during online stage

The procedure of online stage is summarized as follows:

Step 1: Assume that a contaminated EEG signal is represented as $E_{D \times N}$, where D and N stand for the signal length and number of channels, respectively. For computation, the $E_{D \times N}$ is concatenated into a vector in R^T , with $T = D \times N$.

Step 2: Divide the R^T into c segments, which each segment is recorded as Seg_τ , where $\tau \in N_+$, $\tau \leq c$ denotes the sequence number of segments. And the Seg_τ is an n dimensional vector, which corresponds to n inputs of DLN.

Step 3: Standardize the Seg_τ according to the Eq. (6).

Step 4: Input the Seg_τ to the trained DLN and obtain the output data.

Step 5: Obtain the corrected EEG according to the Eq. (7).

2.2. Simple introduction of compared methods

The proposed method is compared with four methods in this paper, including the shallow network SAE, ICA, kurtosis-ICA (K-ICA) and Second-order blind identification (SOBI).

SAE is a kind of a shallow network, which can be used to learn high-order statistical moments. However, its learning ability is limited compared with the DLN. A detailed introduction of the SAE can be seen in [47].

ICA is effective in multi-type artifacts removal and it should be a good benchmark method for educational purpose in EEG-related research [8]. The principle of ICA has been introduced in the introduction part.

K-ICA is an improved method of ICA. This method uses kurtosis to determine the threshold to separate the artifacts-affected ICA components from the unaffected components [24]. The most important advantage of using this method is that it can remove OAs automatically. However, some important information may be compromised while removing the OAs.

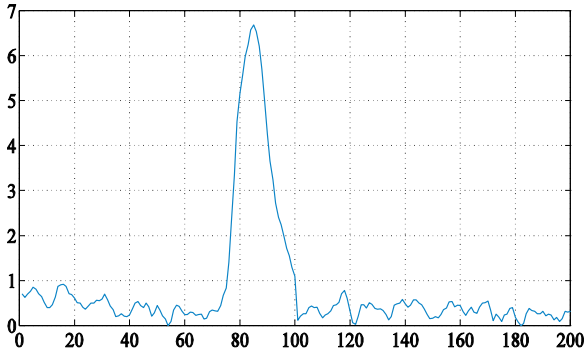
The Second-order blind identification (SOBI) algorithm utilize the difference of source signal's space and frequency. Firstly, it calculates the whitening data from multiple time delays to get a set of correlation coefficient matrices. Secondly, a set of correlation coefficient matrices are processed with optimal joint diagonalization. Finally, the correlations of independent sources are eliminated from source signal. SOBI [49] is an emerging signal processing technique that can be used to facilitate source analysis from EEG [50] and MEG data [51]. SOBI works to minimize the relatedness among the final recovered source signals that make up the mixtures of signals.

2.3. Performance metrics

Three metrics are used to assess the performances of the proposed method, including power spectral density (PSD), root mean square error (RMSE) and EEG classification accuracy.

PSD defines how the power of the signal or time series varies with frequency. This paper uses PSD to evaluate the OAs removal effect.

RMSE is "Root Mean Square Error". It's always to evaluate the precision of the measuring value and the true value. In this paper, RMSE is used to evaluate the reconstruction ability of DLN and the OAs removal ability of the proposed method. It is the meaning of

Fig. 7. An EEG_{std} for 2s.

reconstruction ability of DLN. For the reconstruction ability of DLN, the RMSE is seen as a reconstruction error and defined as follows:

$$RMSE = \sqrt{\frac{1}{n} \sum_{k=1}^n (EEG_{otp}(k) - EEG_{inp}(k))^2} \quad (8)$$

where EEG_{inp} denotes the input data of the DLN. For the OAs removal ability of the proposed method, the RMSE is defined as follows:

$$RMSE = \sqrt{\frac{1}{n} \sum_{k=1}^n (EEG_{crt}(k) - EEG_{NOE}(k))^2} \quad (9)$$

where EEG_{NOE} denotes an NOE for length n . The value of RMSE shows the difference between the corrected and original non-contaminated signal. It depicts the degree of the two-signals' approximation. In this case, with the small RMSE value, it could make conclusion that the output EEG reconstructed signal has contained more of the effective original EEG component.

The proposed method is used on the motor imagery EEG data with its related frequency domains range under 40 Hz, so the motor imagery classification accuracy is a good way to test whether this method would cause the loss of useful information.

3. Data description

The dataset used in this paper is from “Data sets 1” for BCI Competition IV, which is launched on July 3rd 2008. This data set consists of EEG data from 7 subjects of a study published in [48].

However, data for subjects 3–5 have been artificially generated. For each subject, two classes of motor imagery were selected from the three classes: left hand, right hand, and foot (side chosen by the subject; optionally also both feet). Each subject has 200 trails of motor imagery and each trail lasts for more than 6 s. EEG signals were recorded from 59 channels, which positioned over sensorimotor areas densely. Signals were band-pass filtered between 0.05 and 200 Hz and then digitized at 100 Hz with 16bit ($0.1 \mu V$) accuracy. More details were described in [48]. In this paper, we selected Subject 1, 2, 6 and 7 to test the proposed method, because their EEG data are real. In addition, EEG data from three individuals tested in our lab are added to check the validity of the proposed method. Neuracle WirelssEEG Acquisition System is used to acquire the signals with 32 channels EEG. The acquisition system has high precision(24-bit) with low input noise($<0.4 \mu V_{rms}$) sampling technology to guarantee the signal quality while recording small amplitude signal. Each individual also has 200 trials of motor imagery and each trial also lasts for more than 6 s.

4. Experiments and results

4.1. DLN training

In this paper, n equals 100 input units are inputted into DLN. For each subject, the EEG data is divided into two parts (Part 1 and Part 2), while each part consists of 100 trails. Part 1 is used to train and test the DLN, and Part 2 is used to remove OAs and test the proposed method. Take Subject 1 for example, according to Section 2.2, we obtain 16520 training samples to train the DLN and 15458 test samples to test the reconstruction ability of the trained DLN from Part 1. Through a large number of experiments, we find that there are mainly two factors affecting the reconstruction ability of DLN. One is the sparsity parameter ρ (applied in the hidden layers) and the other is the weight of the sparsity penalty term β (details can be seen in [45,47]). Take DLN with a structure of 100-80-100-80-100 for example, the relationship between the reconstruction error (RMSE) and the two factors can be seen in Fig. 8. From Fig. 8 we can see that the RMSE obtains the minimum when $\rho = 0.4$, $\beta = 0.00005$. Thus, $\rho = 0.4$, $\beta = 0.00005$ is taken in this paper. It is noted that an acceptable RMSE reconstruction value could be least for the priority because the least RMSE value shows that the reconstruction EEG segments contain enough effective original EEG segments.

After training, the DLN seems to have “learned” these waveforms adequately enough to be able to reconstruct them. In order

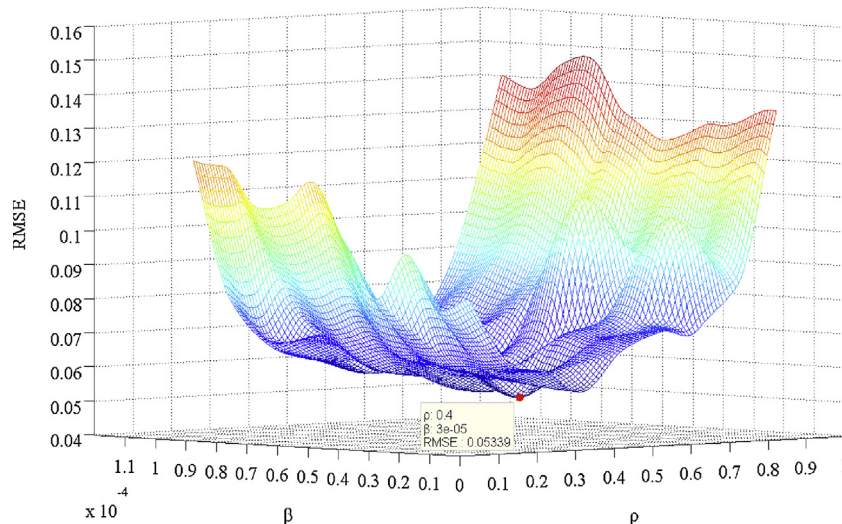


Fig. 8. RMSE's Relationship with sparsity parameter and sparsity penalty.

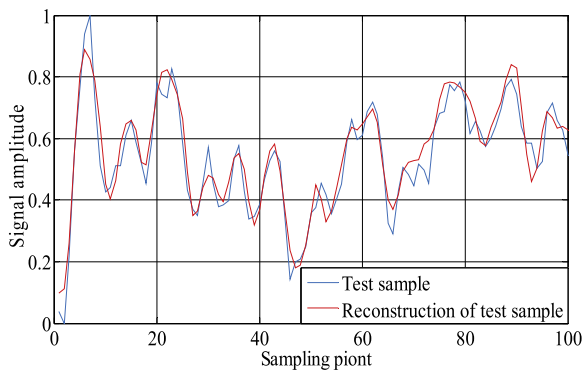


Fig. 9. Reconstruction of test sample by DLN.

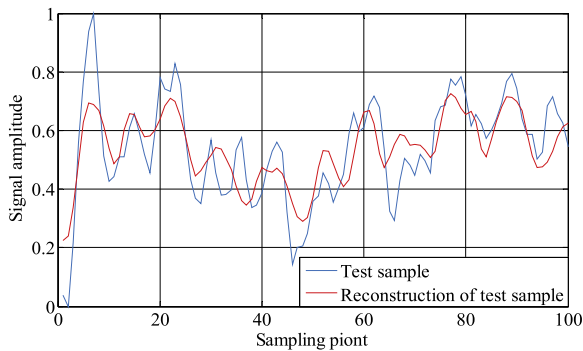


Fig. 10. Reconstruction of test sample by SAE.

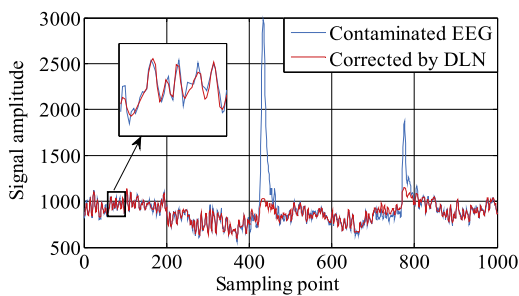


Fig. 11. Removal effect by using DLN.

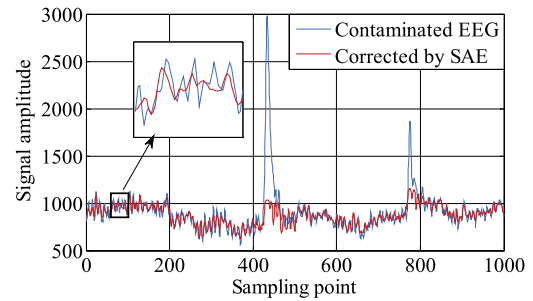


Fig. 12. Removal effect by using SAE.

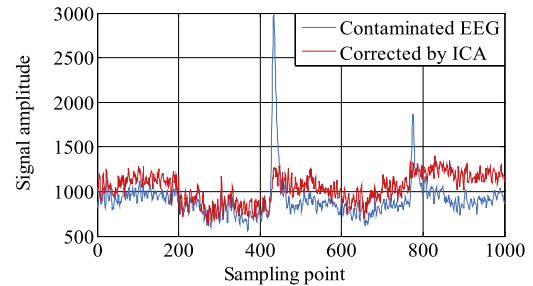


Fig. 13. Removal effect by using ICA.

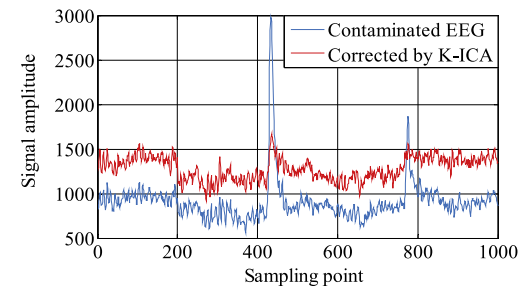


Fig. 14. Removal effect by using K-ICA.

Table 3

RMSE values of corrected EEG signals.

Methods	DLN	SAE	ICA	K-ICA	SOBI
RMSE	1.2838	1.2602	1.4934	1.4260	1.371

to show stronger learning ability of DLN, we have compared the reconstruction ability between DLN and SAE. Figs. 9 and 10 depict reconstructions of one of testing samples drawn by the output of DLN for 3 hidden layers and SAE, respectively. Clearly, the DLN performs better in the details of the signal.

4.2. OAs removal

The OAs removal is processed during the online stage. In this paper, EEG data from Part 2 is used to remove OAs. Take a contaminated EEG segment from Subject 1 for example, Fig. 11 depicts the removal effect by using the proposed method. Fig. 11 shows that the OAs in the contaminated EEG are successfully removed. We also use other four methods (SAE, ICA, K-ICA, SOBI) to make the comparison. Figs. 12–15 show the removal effects by using SAE, ICA, K-ICA and SOBI, respectively. Although, the SAE obtains the similar result as the DLN, the DLN performs better in reconstructing the details of the signal. It is deserved to be mentioned that the reconstruction result on gamma band could not be perfectly coincided, the gamma band contains few information about motor imagery signal for the physiological feature morphology. The EEG band for motor

imagery analysis is concentrate under 40 Hz with the hidden useful information. The tuning ρ and β parameters for DLN reconstruction result could beyond the satisfactory between [0,40] Hz frequency. This paper concentrates on ocular artifact denoising method used for motor imagery rehabilitation system. The reconstruction of DLN on gamma band would influence little to the latter analysis on EEG motor imagery analysis. But if the gamma band could be approximated better, there is no doubt to further optimize performance of the motor imagery rehabilitation system. Fig. 13 shows that the OAs are removed, but the corrected EEG is deformed compared with the original EEG. Fig. 14 shows that the OAs removal is not complete. Moreover, the corrected EEG is moved compared with the original EEG. Fig. 15 shows that the OAs are removed on the whole, but the details have a little deforming effect. Table 3 gives the RMSE of the corrected EEG, which shows that the DLN and SAE method perform better than SOBI, ICA and K-ICA.

In the frequency domain, the proposed method also obtains good result. Fig. 16 shows PSD of the contaminated EEG and the corrected EEG processed by DLN, SAE, ICA, K-ICA and SOBI, respectively. Theoretically, PSD value of an ideal corrected EEG should be

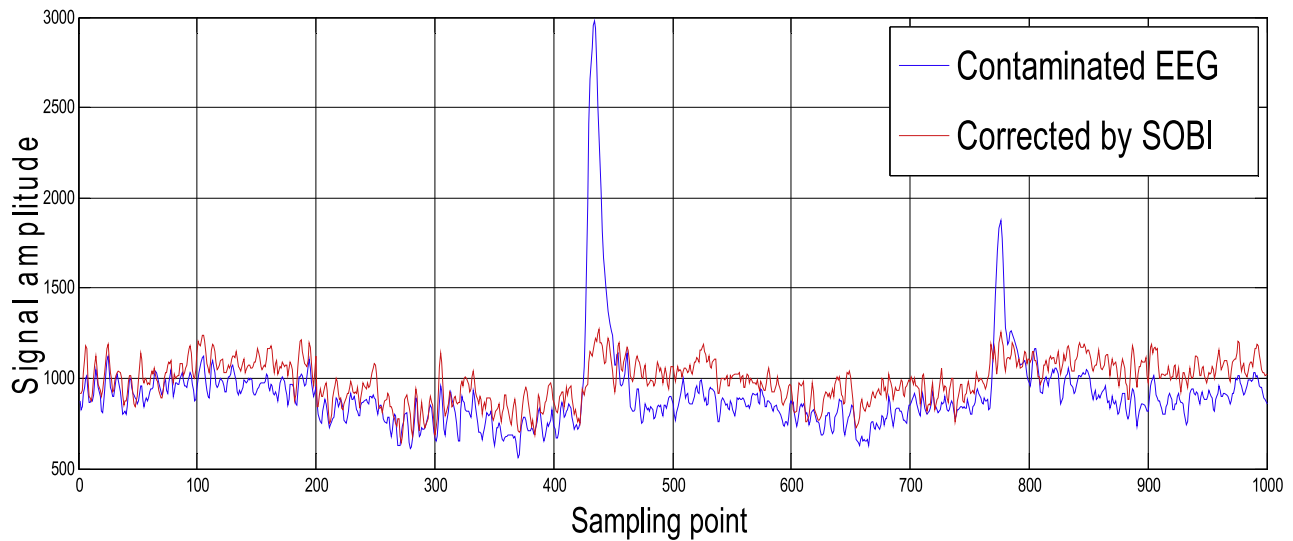


Fig. 15. Removal effect by using SOBI.

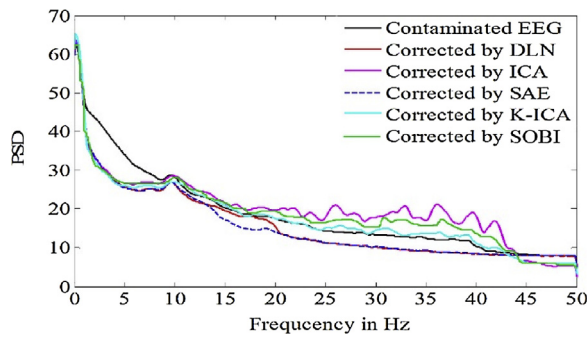


Fig. 16. PSD of contaminated EEG and corrected EEG.

Table 4
Classification accuracy of different methods.

	Contaminated EEG	DLN	SAE	ICA	K-ICA	SOBI
Subject 1	78.2%	79.8%	76.4%	74.6%	78%	77%
Subject 2	76.6%	79.2%	76.2%	71.8%	75.4%	71%
Subject 6	63.6%	66%	63%	62.6%	61.6%	67%
Subject 7	71.4%	75.4%	72.4%	70.2%	71.4%	73%

small in the concentrated frequency ranges of OAs, which should approximate the contaminated EEG as accuracy as possible in other frequency ranges. From Fig. 16 we can see that in the concentrated frequency ranges of OAs (δ , θ and α), PSD values of five corrected EEG signals are all declining. And DLN and SAE perform slightly better than ICA, K-ICA and SOBI, but in the 13–20 Hz area, where many useful EEG information is contained, the DLN method performs better than SAE.

In this paper, the classification accuracy of motor imagery is used as a metric to test whether five methods caused loss of useful information. We use Part 2 from four subjects to calculate the classification accuracies. For each subject, Part 2 contains 100 trials (100 labels). Intercept first 5 s of EEG data for each trial, and then evenly divide these 5-s EEG data into 5 segments. Therefore, we can obtain 500 labels in total and each label corresponds to 1 s of EEG data. Here, the common spatial pattern (CSP) is used for feature extraction and the support vector machine (SVM) for classification. Table 4 shows the classification accuracies of the contaminated EEG and the corrected EEG for 4 subjects. From Table 4

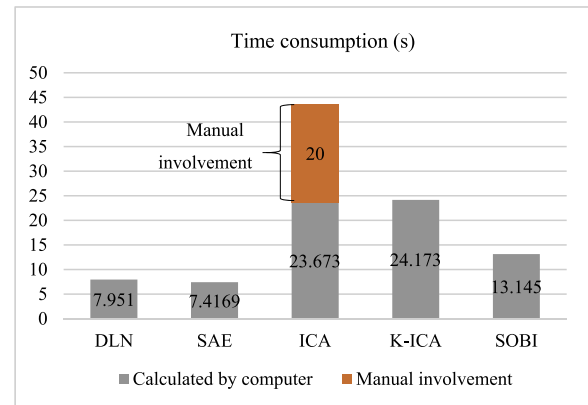


Fig. 17. Computational time of different methods.

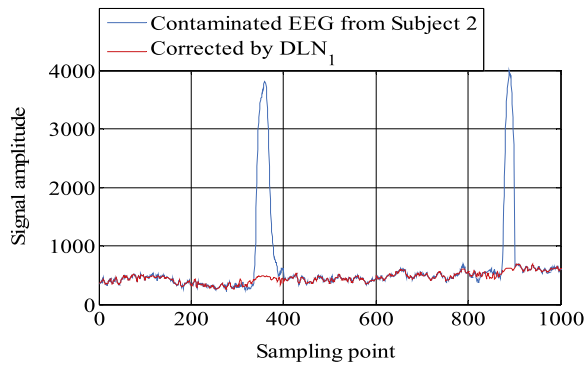
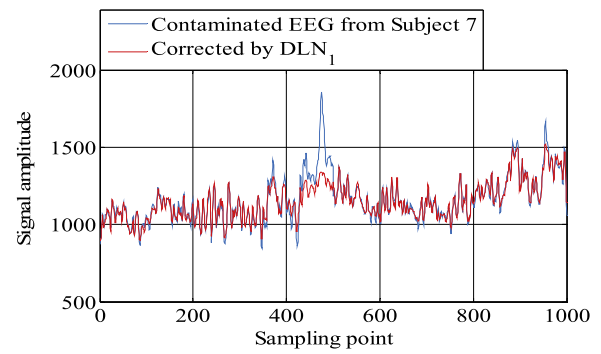
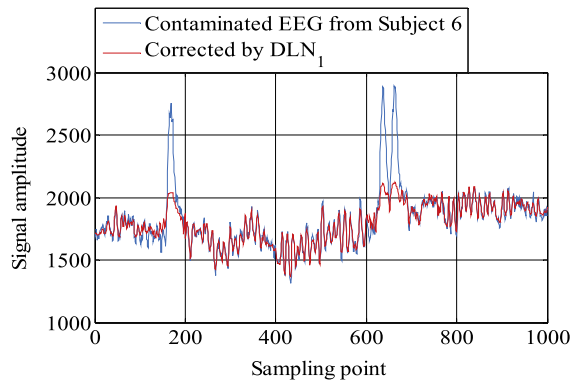
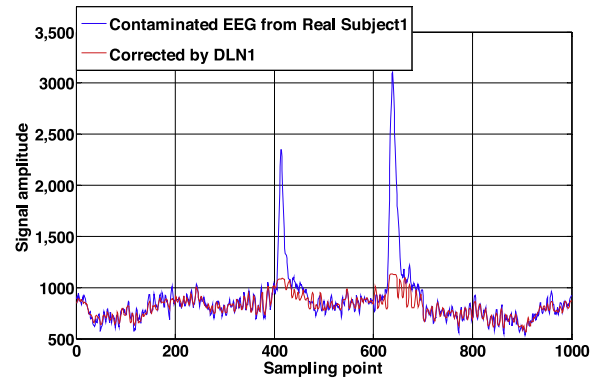
we can see most of classification accuracies processed by four compared methods are reduced in varying degrees, which implies that the four compared methods may cause the loss of useful EEG information.

4.3. Time consumption

Time consumption is a very important issue in OAs removal. The processing time for 200 trials (each trail contains 59 EEG channels and each channel in each trail contains 5 s of EEG data) under different methods is given in Fig. 17. The time is calculated with MATLAB (release R2013a) running on Windows 10 (Intel Core i3-4130 CPU 3.4 GHz processor, 10 GB RAM). Fig. 17 shows that the DLN is much more time saving compared with ICA, K-ICA and SOBI. Because of the fact that ICA needs manual involvement to eliminate OAs, its time consumption fluctuates between 20 and 45 s.

4.4. Generalization ability

The proposed method is proven to have strong generalization ability, which is good for widely used. In this paper, we use the cross-subject testing scenarios to test the generalization ability of the proposed method. For example, we first use the training samples from Subject 1 to train an DLN, and then use this trained DLN to remove OAs from Subject 2, Subject 6, and Subject 7, respec-

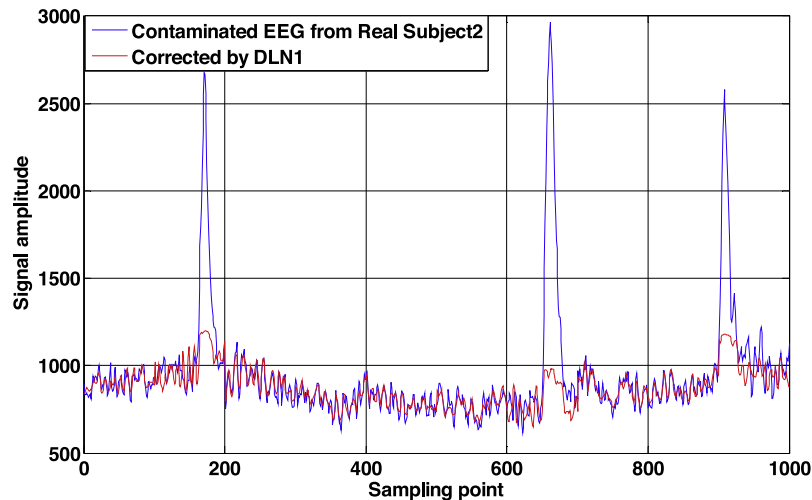
Fig. 18. OAs removal by using DLN_1 for Subject 2.Fig. 20. OAs removal by using DLN_1 for Subject 7.Fig. 19. OAs removal by using DLN_1 for Subject 6.Fig. 21. OAs removal by using DLN_1 for Participant 1.

tively. In addition, three healthy participants are engaged in the experiment for further real validity. For convenience, DLN_i denotes an DLN whose training samples are from Subject i ($i = 1, 2, 6, 7$). Figs. 18–20 give the OAs removal effects by using DLN_1 for Subject 2, Subject 6, and Subject 7, respectively. Figs. 21–23 give the OAs removal effects by using DLN_1 for Participant1, Participant2 and Participant3. From the six figures we can see that OAs are still successfully removed. Table 5 gives comparisons of classification accuracies for Subject 2, Subject 6, and Subject 7 by using DLN_1 and their own DLN. And of course, Table 5 also gives the comparisons of classification accuracies for Participant 1, Participant 2, and Participant 3 by using DLN_1 and their own DLN. As a result, Table 5 has strong evidence to prove that the proposed method has strong

Table 5

Comparisons of classification accuracies by using DLN_1 and their own DLN.

	Contaminated EEG	DLN_1	DLN_2
Subject 2	76.6%	79.6%	79.2%
	Contaminated EEG	DLN_1	DLN_6
Subject 6	63.6%	65.2%	66%
	Contaminated EEG	DLN_1	DLN_7
Subject 7	71.4%	73%	75.4%
	Contaminated EEG	DLN_1	DLN_{p1}
Participant 1	72.5%	73.9%	75%
	Contaminated EEG	DLN_1	DLN_{p2}
Participant 2	75.3%	77.8%	76.9%
	Contaminated EEG	DLN_1	DLN_{p3}
Participant 3	69.5%	71.4%	73.5%

Fig. 22. OAs removal by using DLN_1 for Participant 2.

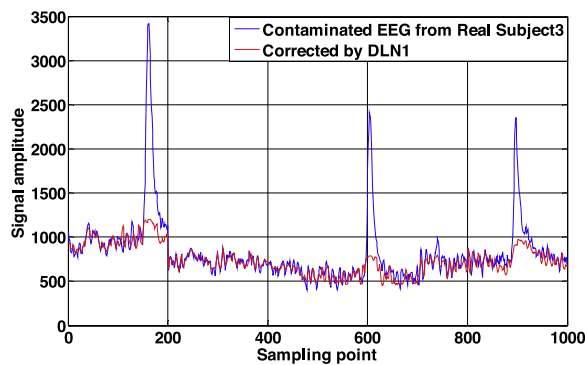


Fig. 23. OAs removal by using DLN₁ for Participant 3.

generalization ability. It should be pointed that we obtained the similar results by using DLN₂, DLN₆, and DLN₇ to the rest subjects including real healthy subjects, respectively.

5. Conclusion and discussion

In this paper, a novel method based on DLN is proposed for removing OAs from contaminated EEG. The proposed method skillfully takes advantage of the structure characteristics and the strong learning ability of DLN, which can remove OAs effectively. Compared with the classical OAs removal methods, the proposed method has many highlights. For example, it doesn't need additional EOG recording in real time and there is no limit to the number of channels. Moreover, the proposed method is much more time saving and it has strong generalization ability, which is useful for online OAs removal.

In the future work, we are going to improve the training method of DLN or try replacing the SAE with other neural networks such as convolutional neural networks (CNN) to strengthen its fitting ability for the details of EEG. And also, more modern denoising algorithms with bigger database should be compared to find the robust ability of this method. Additionally, these networks can also be introduced to other applications such as EEG features extraction and classification.

Acknowledgment

This project is supported by National Natural Science Foundation of China (60975079, 31100709) and the Shanghai Pujiang Program, China (No. 14PJ1431300).

References

- [1] M.M. Mannan, M.Y. Jeong, M.A. Kamran, Hybrid ICA-regression: automatic identification and removal of ocular artifacts from electroencephalographic signals, *Front. Hum. Neurosci.* 10 (2016) 193–209.
- [2] M. De Vos, M. Kroesen, R. Emkes, P300 speller BCI with a mobile EEG system: comparison to a traditional amplifier, *J. Neural Eng.* 11 (2014) 277–286.
- [3] M.A. Kamran, K.S. Hong, Linear parameter-varying model and adaptive filtering technique for detecting neuronal activities: an fNIRS study, *J. Neural Eng.* 10 (2013) 056002.
- [4] L.J. Hoon, L.S. Min, B.H. Jin, CNT/PDMS-based canal-typed ear electrodes for inconspicuous EEG recording, *J. Neural Eng.* 11 (2014) 046014.
- [5] J.E. Kline, H.J. Huang, K.L. Snyder, Isolating gait-related movement artifacts in electroencephalography during human walking, *J. Neural Eng.* 12 (2015) 046022.
- [6] R. Sameni, C. Gouy-Pailler, An iterative subspace denoising algorithm for removing electroencephalogram ocular artifacts, *J. Neurosci. Methods* 225 (2014) 97–105.
- [7] L. Schmäser, A. Sebastian, A. Mobascher, K. Lieb, O. Tüscher, B. Feige, Data-driven analysis of simultaneous EEG/fMRI using an ICA approach *Front. Neuroscience* 8 (2014) 175.
- [8] H.A.T. Nguyen, J. Musson, F. Li, et al., EOG artifact removal using a wavelet neural network, *Neurocomputing* 97 (2012) 374–389.
- [9] N.A. de Beer, W.L. van Meurs, M.B. Grit, et al., Educational simulation of the electroencephalogram (EEG), *Technol. Health Care (Don Mills)* 9 (2001) 237–256.
- [10] L. Sherlin, Diagnosing and treating brain function through the use of low resolution electromagnetic tomography (LORETA), in: *Introduction to Quantitative EEG and Neurofeedback*, Elsevier, New York, 2009.
- [11] J.P. Rosenfeld, A.M. Reinhart, S. Srivastava, The effects of alpha (10-Hz) and beta (22-Hz) “Entrainment” simulation on the alpha and beta EEG bands: individual differences are critical to prediction of effects, *Appl. Psychophys. Biof.* 22 (1997) 1023–1034.
- [12] M.R. Mowla, S.C. Ng, M.S.A. Zilany, et al., Artifacts-matched blind source separation and wavelet transform for multichannel EEG denoising, *Biomed. Signal. Process.* 22 (2015) 111–118.
- [13] N. Mammone, F.L. Foresta, F.C. Morabito, Automatic artifact rejection from multichannel scalp EEG by wavelet ICA, *IEEE Sens. J.* 12 (2012) 533–542.
- [14] D. Wulsin, J. Blanco, R. Mani, et al., Semi-supervised anomaly detection for EEG waveforms using deep belief nets, *International Conference on Machine Learning and Applications* (2010) 436–441.
- [15] T.P. Luu, Y. He, S. Brown, et al., Gait adaptation to visual kinematic perturbations using a real-time closed-loop brain-computer interface to a virtual reality avatar, *J. Neural Eng.* 13 (2016) 036006.
- [16] J. Toppi, M. Riseti, L.R. Quitadamo, et al., Investigating the effects of a sensorimotor rhythm-based BCI training on the cortical activity elicited by mental imagery, *J. Neural Eng.* 11 (2014) 035010.
- [17] Luz María Alonso-Valerdi, F. Sepulveda, Ricardo A. Ramírez-Mendoza, Perception and cognition of cues used in synchronous brain-computer interfaces modify electroencephalographic patterns of control tasks *Front. Hum. Neurosci.* 9 (2015) 636.
- [18] T.P. Jung, S. Makeig, C. Humphries, T. Lee, M. Mckeown, V. Iragui, et al., Removing electroencephalographic artifacts by blind source separation, *Psychophysiology* 37 (2000) 163–178.
- [19] M.A. Klados, C. Papadelis, C. Braun, et al., ICA: a hybrid methodology combining Blind Source Separation and regression techniques for the rejection of ocular artifacts, *Biomed. Signal. Process.* 6 (2011) 291–300.
- [20] S.M. Haas, M.G. Frei, I. Osorio, et al., EEG ocular artifact removal through ARMAX model system identification using extended least squares, *Commun. Inform. Syst.* 3 (2003) 19–40.
- [21] M. Fatourechi, A. Bashashati, R.K. Ward, et al., EMG and EOG artifacts in brain computer interface systems: a survey, *Clin. Neurophysiol.* 118 (2010) 480–494.
- [22] D. Hagemann, E. Naumann, The effects of ocular artifacts on (lateralized) broadband power in the EEG, *Clin. Neurophysiol.* 112 (2001) 215–231.
- [23] D.A. Pizzagalli, Electroencephalography and high-density electrophysiological source localization, in: *Handbook of Psychophysiology*, Cambridge University, New York, 2007, pp. 56–84.
- [24] K.A. Islam, G.V. Tcheslavski, Independent component analysis for EOG artifacts minimization of EEG signals using kurtosis as a threshold, in: *Electrical Information and Communication Technology, IEEE*, 2015, pp. 137–142.
- [25] A. Turnip, E. Junaidi, Removal artifacts from EEG signal using independent component analysis and principal component analysis, in: *International Conference on Technology, Informatics, Management, Engineering, and Environment*, IEEE, 2014, pp. 296–302.
- [26] Hu J, C.S. Wang, et al., Removal of EOG and EMG artifacts from EEG using combination of functional link neural network and adaptive neural fuzzy inference system, *Neurocomputing* 151 (2015) 278–287.
- [27] A. Kilicarslan, R.G. Grossman, J.L. Contreras-Vidal, A robust adaptive denoising framework for real-time artifact removal in scalp EEG measurements, *J. Neural Eng.* 13 (2016) 026013.
- [28] S. Seifzadeh, F. Karim, M. Amiri, Comparison of different linear filter design methods for handling ocular artifacts in brain computer interface system, *J. Comput. Robot.* 7 (2014) 51–56.
- [29] B. Noureddin, P.D. Lawrence, G.E. Birch, Time-frequency analysis of eye blinks and saccades in EOG for EEG artifact removal, *International IEEE/EMBS Conference on Neural Engineering* (2007) 143–164.
- [30] B.H. Yang, L.F. He, L. Lin, Fast removal of ocular artifacts from electroencephalogram signals using spatial constraint independent component analysis based recursive least squares in brain-computer interface, *Front. Inform. Technol. Electron. Eng.* 16 (2015) 486–496.
- [31] D.P. He, G. Wilson, C. Russell, Removal of ocular artifacts from electroencephalogram by adaptive filtering, *Med. Biol. Eng. Comput.* 42 (2004) 407–412.
- [32] H. Ghandeharion, A. Erfanian, A fully automatic ocular artifact suppression from EEG data using higher order statistics: improved performance by wavelet analysis, *Med. Eng. Phys.* 32 (2010) 720–729.
- [33] W.Y. Hsu, C.H. Lin, H.J. Hsu, et al., Wavelet-based envelope features with automatic EOG artifact removal: application to single-trial EEG data, *Expert. Syst. Appl.* 39 (2012) 2743–2749.
- [34] V. Krishnaveni, S. Jayaraman, L. Anitha, K. Ramadoss, Removal of ocular artifacts from EEG using adaptive thresholding of wavelet coefficients, *J. Neural Eng.* 3 (2006) 338–346.
- [35] K.P. Indiradevi, E. Elias, P.S. Sathidevi, S.D. Nayak, K. Radhakrishnan, A multilevel wavelet approach for automatic detection of epileptic spikes in the electroencephalogram, *J. Comput. Biol. Med.* 38 (2008) 805–816.
- [36] M.D. Vos, W. Deburchgraeve, et al., Automated artifact removal as preprocessing refines neonatal seizure detection, *Clin. Neurophysiol.* 122 (2011) 2345–2354.

- [37] I. Winkler, S. Brandl, F. Horn, et al., Robust artifactual independent component classification for BCI practitioners, *J. Neural Eng.* 11 (2014) 035013.
- [38] B. Hazra, A. Roffel, S. Narasimhan, et al., Modified cross-correlation method for the blind identification of structures, *J. Eng. Mech.* 136 (2010) 889–897.
- [39] N.P. Castellanos, A.V. Makarov, Recovering EEG brain signals: artifact suppression with wavelet enhanced independent component analysis, *J. Neurosci. Methods* 158 (2006) 300–312.
- [40] S. Ahmed, L.M. Merino, Z. Mao, et al., A deep learning method for classification of images RSVP events with EEG data, 2013 IEEE Global Conference on Signal and Information Processing (GlobalSIP) (2013) 33–36.
- [41] J.F. Saa, M. Çetin, ÇA latent discriminative model-based approach for classification of imaginary motor tasks from EEG data, *J. Neural Eng.* 9 (2012) 26020–26028.
- [42] D.F. Wulsin, J.R. Gupta, R. Mani, et al., Modeling electroencephalography waveforms with semi-supervised deep belief nets: fast classification and anomaly measurement, *J. Neural Eng.* 8 (2011) 53–57.
- [43] G.E. Hinton, S. Osindero, Y.W. Teh, A fast learning algorithm for deep belief nets, *Neural. Comput.* 18 (2006) 1527–1554.
- [44] B. Schölkopf, J. Platt, T. Hofmann, Greedy Layer-wise Training of Deep Networks, MIT Press, 2007, pp. 153–160.
- [45] A. Ng, 2011. Stacked Autoencoders http://ufldl.stanford.edu/wiki/index.php/Stacked_Autoencoders.
- [46] D. Rueda-Plata, R. Ramos-Pollán, F.A. González, Supervised greedy layer-wise training for deep convolutional networks with small datasets, in: Computational Collective Intelligence, Springer International Publishing, 2015, pp. 275–284.
- [47] A. Ng, 2011. Sparse autoencoder. CS294A Lecture notes. 72, 1–19.
- [48] B. Blankertz, G. Dornhege, M. Krauledat, et al., The non-invasive Berlin Brain-Computer Interface: fast acquisition of effective performance in untrained subjects, *Neuroimage* 37 (2007) 539–550.
- [49] A. Belouchrani, K. Abed-Meraim, J.-F. Cardoso, E. Moulines, A blind source separation technique using second-order statistics, *IEEE Trans. Signal Process.* 5 (1997) (1997) 434–444.
- [50] C.A. Joyce, I.F. Gorodnitsky, M. Kutas, Automatic removal of eye movement and blink artifacts from EEG data using blind component separation, *Psychophysiology* 41 (2004) 313–325.
- [51] A.C. Tang, B.A. Pearlmutter, M. Zibulevsky, S.A. Carter, Blind source separation of multichannel neuromagnetic responses, *Neurocomputing* 32 (2000) (2000) 1115–1120.