

뇌심부자극의 딥러닝 기반 아티팩트 제거 기법 개발

성채현, 류경찬, 최혜진, 차미나, 한주연, 김진모*, 최지웅*
대구경북과학기술원

seongchae@dgist.ac.kr, ryukyungchan22@dgist.ac.kr, ch.hyejin@dgist.ac.kr, mina@dgist.ac.kr,
juyeon.han@dgist.ac.kr, *jmkim@@dgist.ac.kr, *jwchoi.dgist@gmail.com

Development of Deep Learning-Based Artifact Removal Method for Deep Brain Stimulation

Chae Hyeon Seong, Kyung Chan Ryu, Hye Jin Choi, Mi Na Cha, Ju Yeon Han,
Jinmo Kim*, Ji-Woong Choi*
Daegu Gyeongbuk Institute of Science and Technology (DGIST)

요 약

뇌심부자극 (Deep Brain Stimulation, DBS)은 파킨슨병 치료법 중 하나로, DBS 와 동시에 뇌신호를 기록할 경우 자극 아티팩트로 인해 신호가 오염되어 기록되는 문제가 발생한다. 이에 오염된 신호를 복원하는 과정이 필요하며, 이를 위해 효과적인 자극 아티팩트 제거 (Stimulation Artifact Cancellation, SAC) 기술이 요구된다. 기존의 선형 기반 SAC 기법은 아티팩트가 비사인모양 파형을 갖거나 시간에 따라 달라지는 파형을 가질 경우 아티팩트를 제거하는 데 한계를 가진다. 본 기법에서는 이를 해결하기 위해 입력 신호의 다중 해상도 분석 (Multi-Resolution Analysis, MRA)을 학습 데이터로 사용하는 Convolutional Neural Network (CNN) 구조 기반의 SAC 기법을 제안한다. 제안 기법의 효과를 평가하기 위해 시뮬레이션을 통해 데이터를 생성하였으며, 결과적으로 제안 기법은 시간 및 주파수 영역의 정보를 모두 학습함으로써 자극 아티팩트의 세기가 달라지는 상황에서도 기존 알고리즘 대비 원래의 뇌신호를 거의 손상 없이 복원하는 성능을 나타낸다. 본 제안 기법은 뇌신호를 기록 및 자극하는 각종 뇌-컴퓨터 인터페이스 발전에 기여할 것으로 기대된다.

I. 서 론

초고령화 시대의 도래와 함께 파킨슨병과 같은 퇴행성 뇌질환의 유병률이 증가하고 있다. 파킨슨병의 대표적인 증상 완화 방법 중 하나인 DBS 는 기저핵 부위에 다중 채널 전극을 삽입하여 전기 자극을 제공하는 방법으로, 자극과 동시에 자극에 대한 반응을 관찰하기 위해 자극과 뇌신호의 기록이 동일한 전극의 인접 채널에서 동시에 이루어진다. 이 과정에서 발생하는 자극 아티팩트로 인해 뇌신호가 오염되는 문제가 나타난다. 뇌신호의 정확한 분석을 위해서는 오염된 신호를 복원하는 것이 필수적이며, 이를 위해 효과적인 SAC 기술이 요구된다.

기존의 SAC 방식, 예를 들어 Notch 필터와 같은 전통적인 필터링 방법 [1]이나 PARRM (Period-based Artifact Reconstruction and Removal Method) [2]과 같은 최신의 SAC 알고리즘은 자극 아티팩트가 비사인모양 파형을 갖거나 시간에 따라 달라지는 모양을 가질 경우, 이를 완전히 제거하지 못하거나 일부 뇌신호를 손상시킨다는 한계점이 존재한다. 본 연구에서는 이러한 한계점을 극복하기 위해 딥러닝 기반의 새로운 SAC 알고리즘을 제시한다.

II. 본론

본 논문에서는 시뮬레이션을 통해 자극 아티팩트에 의해 오염된 뇌신호와 자극이 없는 상태인 ground truth (GT) 뇌신호를 생성하였으며, 이를 학습하여 자극 아티팩트가 제거된 뇌신호를 출력하는 딥러닝 모델을 개발하였다. 다양한 실험을 통해 backbone 구조, 입력

모달리티, 손실 함수를 변화시켜 최적의 모델을 탐색하였으며, 기존 SAC 기법과의 성능을 비교분석 하였다.

i. 데이터 생성 방법

a. 뇌신호 시뮬레이션 데이터 생성

Python 의 NeuroDSP 모듈을 이용하여 자극이 없는 상태에서 2 초 동안 2000 Hz 의 샘플링 주파수로 수집된 신경세포의 시냅스 활동과 특정 주파수의 진동을 고려하여 10000 개의 뇌신호 시뮬레이션 데이터를 생성하였다.

b. 시뮬레이션 기반 자극 아티팩트 생성

자극 아티팩트는 실제 뇌 속에서 DBS 를 수행할 때 자극 파형이 비대칭적으로 변하는 것을 반영하여 125 Hz 자극으로 생성하였으며 평균이 0, 표준편차가 0.3 인 가우시안 노이즈를 더해줌으로써 데이터 수집 시 포함되는 노이즈를 고려하였다. 모델 학습에 활용할 오염 신호는 뇌신호 시뮬레이션 데이터와 자극 아티팩트 시뮬레이션 데이터를 합성하여 생성하였다. 자극이 필요할 때만 DBS 를 수행하고, 효과에 따라 자극 세기를 변경하는 상황을 고려하여 0.2초부터 1.8초까지 1.6초간 자극하고, 1초까지는 GT 뇌신호의 10 배, 이후엔 15 배로 자극 세기를 변경하였다.

ii. 딥러닝 모델 실험 환경

딥러닝 모델 구현과 학습은 PyTorch 프레임워크를 사용하여 수행하였다. 학습 과정에서 하이퍼파라미터 설정을 위해 grid search 를 수행하였으며, 모든 모델에 대해 epoch 수는 500, optimizer 는 Adam, 손실 함수는 평균제곱오차 (Mean Squared Error, MSE), 활성화 함수

표 1. 다양한 모델 구조로 실험한 결과 MSE.

Input	Loss	Backbone structure		
		MLP	CNN	LSTM
시계열	L_t	133.8 ± 22.19	6.501 ± 1.487	19.15 ± 2.768
	$L_t + L_{psl}$	134.0 ± 21.82	7.096 ± 1.690	28.47 ± 4.729
시계열 + PSD	$L_t + L_{psl}$	5.675 ± 1.445	12.19 ± 2.350	2.457 ± 0.5905
MRA	L_{MRA}	9.381 ± 1.517	2.160 ± 0.4418	2.306 ± 0.4508

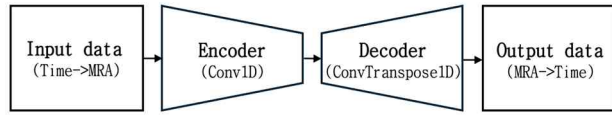


그림 1. 채택한 1D CNN 모델 구조. 각 레이어의 in_feature 은 SWT 를 통해 얻은 MRA 의 각 레벨로 설정하였다. (Encoder: in_feature=6 out_feature=512, GELU, MaxPooling (pool size=(2,2)), Decoder: in_feature=512 out_feature=6).

는 GELU, batch size 는 32, 학습-테스트 데이터셋 분할 비율은 8:2 로 설정하였다. Input 및 hidden size 와 layer 수, 학습률은 모델별로 최적화하였다.

iii. 베이스라인 구현

a. Notch filter

Python SciPy 모듈의 iirnotch 와 filtfilt 함수를 이용하여 125 Hz 의 배수로 나타나는 고조파 성분을 제거하기 위한 Notch filter 를 설계하였다. Filter 의 대역폭을 제어하는 Q-factor 는 30 으로 설정하였다.

b. PARRM

PARRM 필터 설계 시 아티팩트 주파수 값으로 125 Hz 를 제공하였으며, 윈도우 크기를 결정하는 window_time 은 0.99, 뇌신호 훼손을 방지하기 위한 skip_time 은 window_time / 30, 시간 영역에서 데이터 그룹을 지정하는 per_distance 는 0.01 로 최적화하였다.

iv. SAC 성능 평가

평가지표는 SAC 가 수행된 뇌신호와 GT 뇌신호를 각각 0-100 범위로 min-max scaling 한 뒤 시간 영역의 MSE 와 주파수 영역에서의 MSE 평균값으로 설정하였다.

III. 결론

i. 최적화 SAC 모델 탐색

모델의 backbone 구조로 multilayer perceptron (MLP), CNN, long short-term memory (LSTM)를 사용하였으며, 모델의 입력 모달리티로는 시계열을 학습시키는 경우, 시계열과 PSD 의 특징을 뽑아내어 학습시키는 경우, MRA 를 학습시키는 경우를 비교하였다.

그 결과 (표 1, 그림 1), 오염된 뇌신호와 GT 뇌신호의 시계열 데이터를 SWT (Stationary Wavelet Transform)하여 얻은 MRA 결과를 학습 데이터로 사용하는 1D CNN 모델의 성능이 가장 우수하였다.

MRA 는 시간-주파수 정보를 나타내는데, 특히 뛰어난 시간 해상도를 가진다는 특징이 있다. 본 연구에서 사용된 뇌신호는 시간에 따른 비주기적인 변화가 많기 때문에, 시간 해상도가 높으면서 주파수 정보도 포함하고

표 2. SAC 된 오염 뇌신호와 GT 뇌신호 간 MSE.

	MSE
Notch [1]	1489 $\pm$ 1305
PARRM [2]	4023 $\times 10^2 \pm 1161 \times 10^3$
Ours	2.160 $\pm$ 0.4418

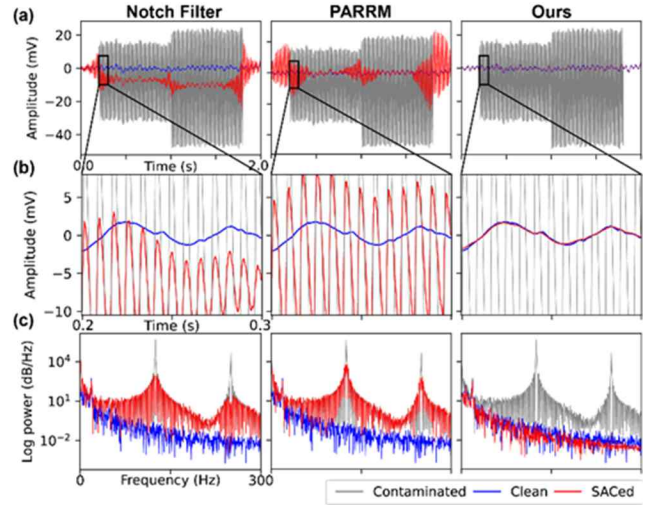


그림 2. 시간 영역과 주파수 영역에서의 SAC 결과 시각화. (a) 시계열 (b) a 를 확대한 결과 (c) PSD.

있는 MRA 를 활용하여 모델을 학습했을 때 모델이 뇌신호의 특성을 잘 학습할 수 있었던 것이라 생각된다.

ii. 베이스라인과의 비교

표 2 와 그림 2 으로부터 Notch filter 와 PARRM 에 비해 제안된 모델의 SAC 성능이 훨씬 우수함을 확인할 수 있었다. 특정 주파수만 제거하는 Notch filter 는 비대칭 자극 아티팩트에 대하여 취약하며 시간에 따라 달라지지 않는 파형을 추정하는 PARRM 은 진폭이 변하는 자극 아티팩트에 대하여 취약하다.

두 가지 기법과 달리 본 논문에서 제안하는 모델은 시간 및 주파수 영역의 정보를 모두 고려하여 학습한다. 따라서 본 모델은 자극 아티팩트의 세기가 달라지는 상황에서도 기존 기법들보다 우수한 성능을 보였으며 원래의 뇌신호를 거의 손상 없이 복원하는 뛰어난 정확도를 나타냈다. 이러한 결론은 뇌신호를 기록하고

자극하는 각종 뇌-컴퓨터 인터페이스 발전 기술에
기여할 것으로 기대한다.

ACKNOWLEDGMENT

본 연구는 과학기술정보통신부에서 지원하는 DGIST
UGRP (Undergraduate Group Research Program)에
의해 수행되었습니다 (과제번호: 2024020011).

참 고 문 헌

- [1] A. de Cheveigné and I. Nelken, "Filters: When, Why, and How (Not) to Use Them," *Neuron*, **102**(2), 280–293, (2019).
<https://doi.org/10.1016/j.neuron.2019.02.039>.
- [2] E. M. Dastin-van Rijn, et al., "Uncovering biomarkers during therapeutic neuromodulation with PARRM: Period-based Artifact Reconstruction and Removal Method," *Cell Reports Methods*, **1**(2), 100010, (2021).
<https://doi.org/10.1016/j.crmeth.2021.100010>.