

Wrapped One-Class Riemannian EEG Classifier for BCI-Based Detection of Anesthetic States

Valérie Marissens Cueva^{1,2†}, Thibault de Surrel^{3†}, Laurent Bougrain^{2,4},
Seyed Javad Bidgoli⁵, Guy Cheron⁶, Ana Maria Cebolla Alvarez⁶, Claude Meistelman⁷,
Fabien Lotte¹, Florian Yger⁸, Sébastien Rimbart¹
† These authors contributed equally to the paper.

¹Inria Center at Univ. Bordeaux / LaBRI, Talence, France; ²Université de Lorraine, CNRS, LORIA, Nancy, France; ³LAMSADE, CNRS, PSL Univ. Paris-Dauphine, France; ⁴Sorbonne Université, ICM, CNRS, Inria, Inserm, Paris, France; ⁵CHU Brugmann, Bruxelles, Belgium; ⁶Laboratory of Neurophysiology and Movement Biomechanics, Université Libre de Bruxelles, Bruxelles, Belgium; ⁷DevAH, University of Lorraine, Nancy, France; ⁸LITIS, INSA Rouen-Normandy, France;
E-mail: thibault.de-surrel@lamsade.dauphine.fr & valerie.marissens@loria.fr

ABSTRACT: Brain-computer interfaces based on ElectroEncephaloGraphy (EEG) often represent efficiently brain signals as covariance matrices and leverage Riemannian geometry for classification of mental states. However, current approaches typically require multiple classes for training. In many applications such as intraoperative monitoring of consciousness during anesthesia, only data from one class, such as the awake state, may be available, requiring one-class classification methods. We propose a novel Riemannian one-class classifier called the One-Class Wrapped Gaussian. Unlike existing methods that rely solely on the Riemannian mean, our approach incorporates second-order statistical information by using an anisotropic Gaussian-like distribution on the manifold of covariance matrices. We validated our method on EEG data from 19 patients undergoing general anesthesia. Results show that our classifier significantly outperforms state-of-the-art one-class methods for distinguishing between awake and anesthetized states, for clinically relevant electrode numbers. We further demonstrate the robustness of our approach by testing configurations with reduced electrode numbers, confirming its feasibility for real-world surgical settings where electrode placement is constrained.

INTRODUCTION

Current Brain-Computer Interfaces (BCIs) based on ElectroEncephaloGraphy (EEG) typically rely on a subject-specific calibration, due to the non-stationarity nature of EEG signals and the large inter-subject variability [1]. Supervised training is generally required to discriminate between mental states, which implies the availability of labeled data for each class. In many real world applications, however, this requirement may be difficult or even impossible to satisfy. In particular, clinical monitoring scenarios often only allow access to data from a single reference condition, while the alternative cannot be labeled in advance. For instance, to monitor the transition between two states of general anesthesia, awake

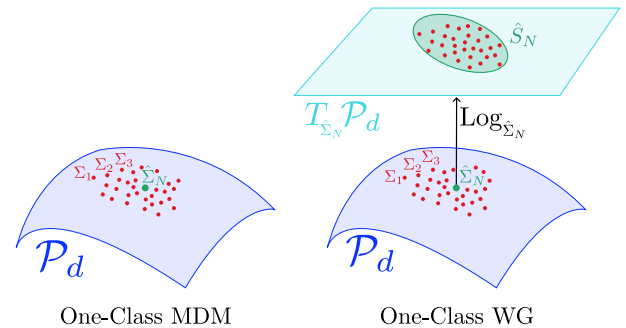


Figure 1: Illustration of the two compared algorithms: *One-Class MDM* (left) and *One-Class WG* (right). While *One-Class MDM* estimates only the Riemannian mean $\hat{\Sigma}_N$ of the training covariance matrices $\Sigma_1, \Sigma_2, \Sigma_3, \dots$, *One-Class WG* also estimates a dispersion matrix in the tangent space $T_{\hat{\Sigma}_N} \mathcal{P}_d$.

or under anesthesia, a classifier would require data from both states. However, before surgery, data can only be obtained from an awake patient. This constraint motivates the use of one-class classification strategies [2, 3], in which a model is trained exclusively on examples from a target class and detects deviations from this reference. While Riemannian geometry-based methods, relying on covariance matrices, have become a standard framework for EEG decoding in BCIs [4], existing approaches usually rely on two-class (or more) training.

Among Riemannian algorithms, the Minimum Distance to the Riemannian Mean (MDM) classifier [5] has been adapted to a one-class framework as a baseline method [6], where the target class is modeled by its geometric centroid. Although simple and efficient, this approach implicitly assumes that covariance matrices are equally distributed in all directions, i.e., that they follow an isotropic distribution [7]. As a result, it cannot explicitly model any anisotropic dispersion of data on the Riemannian manifold of symmetric positive definite (SPD) matrices.

In this work, we introduce a novel one-class classifier based on the Wrapped Gaussian distribution defined on the Riemannian manifold of SPD matrices. The proposed

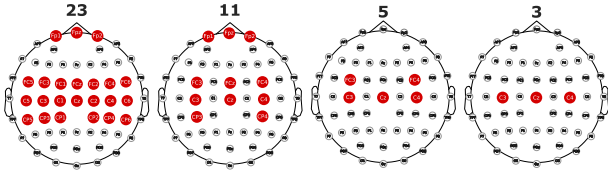


Figure 2: Electrode layouts tested compared to the full 64 electrodes montage.

method provides an explicit -and anisotropic- probabilistic description of the target class preserving the intrinsic geometry of covariance matrices. We evaluate the proposed approach in an EEG-based BCI designed to detect motor cortical activity [8]. Specifically, the motor pattern elicited by median nerve stimulation (MNS) is analyzed, as discriminating between MNS alone and MNS combined with motor intention (MI) has been shown to be more effective than distinguishing MI from rest [9]. The objective is to differentiate EEG responses recorded in awake subjects from those observed under altered brain dynamics induced by general anesthesia. Ultimately, this framework aims to contribute to the detection of intraoperative awareness by identifying potential motor intention from the anesthetized patient. This setting naturally defines a one-class learning problem, as only the reference condition can be collected before the surgery (here EEG during the awake state) for subject-specific calibration, while the alternative condition (here subject-specific EEG during anesthesia) is not accessible beforehand. Beyond this specific use case, the proposed method could be relevant for a broader range of BCI applications involving reference-only training, such as passive BCIs, vigilance monitoring, or brain-state change detection [3].

MATERIALS AND METHODS

Data description: Protocol: This study was conducted at the hospital of Brugmann in Belgium, approved by its ethical committee (CE 2021/225), and registered at EUDRACT (2021-006457-56) and ClinicalTrials.gov (NCT05272202) [8]. A total of 28 subjects scheduled for surgery under general anesthesia were initially enrolled. Four subjects (1, 2, 5 and 28) were excluded either due to technical or logistical reasons beyond our control. Five additional subjects (9, 10, 11, 15 and 25) were excluded to ensure dataset homogeneity in terms of stimulation side and anesthesia protocol. The final cohort consists of 19 subjects (10 females, 47.47 ± 12.51 years old). Loss of consciousness was induced using propofol, while the rest of the anesthesia protocol was determined by the attending anesthesiologist. Ketamine was not administered. The target propofol concentrations were recorded simultaneously with the EEG signals. Neuromuscular blocking agents (Rocuronium) were administered when clinically required (13 patients: 4, 6, 7, 8, 12, 13, 16, 17, 18, 22, 23, 24 and 26).

EEG acquisition and MNS: MNS consisted of 0.2 ms square electrical pulses, generated by the Micromed SD Ltm Stim Energy stimulator, and delivered to the left

wrist through a pair of grass gold cup electrodes. Stimulation intensity was individually adjusted to elicit a visible thumb twitch [10] and remained below 15 mA. EEG signals were recorded with a 64-channel ANT Neuro cap, and the eego mylab system at a sampling rate of 4096 Hz to preserve high-frequency content for future somatosensory evoked potential analyses. In this study, EEG signals were downsampled to 128 Hz and band-pass filtered (8-30 Hz). Artifacts were then visually rejected using EEGLAB[11] (28 ± 37 rejected epochs), and the data were epoched into 750 ms windows (250-1000 ms after MNS). Several electrode montages were selected (see Fig. 2), as current anesthesia monitors use frontal electrodes to detect changes in consciousness associated with general anesthesia [12], whereas, given that this is a motor task evoked by stimulation of the median nerve, the electrodes located in the motor cortex are also of interest.

Recording sessions and training setup: Each run had 150 MNS with a 3-4 s inter-stimulus interval across two sessions. In the preoperative session, at least one run was acquired with the awake patient, yielding 94-499 MNS trials per subject. During the intraoperative session, several EEG runs were recorded, the number of runs depending on the surgery duration, resulting in 415-2390 MNS trials in the anesthetized condition. One-class classifiers were trained on the first 65 preoperative trials, the awake class, and evaluated on the remaining preoperative trials (awake test class) and on the entire intraoperative set (anesthetized test class), resulting in 444-2472 test trials per subject. Classifications were performed using the MNE [13], Scikit-learn [14] and pyRiemann [15] packages in Python 3.10¹.

Riemannian tools used in BCI: In this paper, we represent an EEG trial X by its covariance matrix $\Sigma \in \mathbb{R}^{d \times d}$ where d is the number of EEG electrodes used. This covariance matrix contains information on the link between the different electrodes. Its underlying specific mathematical structure is a SPD matrix. Using the Riemannian geometry of the set $\mathcal{P}_d$ of SPD matrices has lead to state-of-the-art classifiers in EEG based BCIs [4]. Given an EEG trial X , one can compute its covariance matrix $\Sigma = \frac{1}{t-1}XX^T$, where t is the number of sampled time points in each trial. In this paper, we rely on the *Affine-Invariant Riemannian geometry* [16] on $\mathcal{P}_d$. This geometry induces a distance between two covariance matrices Σ_1 and Σ_2 , defined as follows:

$$\delta_R(\Sigma_1, \Sigma_2) = \|\log(\Sigma_1^{-1/2}\Sigma_2\Sigma_1^{-1/2})\|_F,$$

where $\log$ denotes the matrix logarithm and $\|\cdot\|_F$ the Frobenius norm. The tangent space $T_\Sigma\mathcal{P}_d$ is a linear approximation of the curved space $\mathcal{P}_d$ at Σ and can be identified with the flat space of $d \times d$ symmetric matrices $\mathcal{S}_d$, and thus, is of dimension $d(d+1)/2$. To move between $\mathcal{P}_d$ and $T_\Sigma\mathcal{P}_d$, we use the Riemannian exponential Exp_Σ and its inverse, the logarithm map, defined

¹Find the code at https://github.com/valeriemarissens/one_class_wrapped_gaussian

for $\Sigma_1, \Sigma_2 \in \mathcal{P}_d$ and $U \in T_{\Sigma_1} \mathcal{P}_d$ as:

$$\begin{aligned} \text{Exp}_{\Sigma_1}(U) &= \Sigma_1^{1/2} \exp(\Sigma_1^{-1/2} U \Sigma_1^{-1/2}) \Sigma_1^{1/2}, \\ \text{Log}_{\Sigma_1}(\Sigma_2) &= \Sigma_1^{1/2} \log(\Sigma_1^{-1/2} \Sigma_2 \Sigma_1^{-1/2}) \Sigma_1^{1/2}. \end{aligned} \quad (1)$$

Finally, the Riemannian mean of N covariance matrices $\Sigma_1, \dots, \Sigma_N$ is the unique minimizer $\hat{\Sigma}_N$ of

$$\hat{\Sigma}_N = \arg \min_{\Sigma \in \mathcal{P}_d} \sum_{i=1}^N \delta_R^2(\Sigma, \Sigma_i). \quad (2)$$

It can be computed using a Riemannian gradient descent.

Wrapped Gaussian on $\mathcal{P}_d$: Let us now define a probability distribution on the manifold $\mathcal{P}_d$ of SPD matrices. This distribution is a version of a multivariate Gaussian defined directly on $\mathcal{P}_d$. For this, we need to define the vectorization $\text{Vect}_{\Sigma}(U)$ at $\Sigma \in \mathcal{P}_d$ of a symmetric matrix U , that is, taking the upper triangular portion of $\Sigma^{-1/2} U \Sigma^{-1/2}$, with off-diagonal elements scaled by $\sqrt{2}$, forming a vector in $\mathbb{R}^{d(d+1)/2}$. One can also consider the inverse operation, $\text{Unvect}_{\Sigma}(t)$, that takes a vector $t \in \mathbb{R}^{d(d+1)/2}$, and creates a symmetric matrix. Let $\Sigma \in \mathcal{P}_d$ and S be a SPD matrix of size $d(d+1)/2 \times d(d+1)/2$, then we say that the random variable X on $\mathcal{P}_d$ follows a *wrapped Gaussian*, denoted $\text{WG}(\Sigma, S)$ if

$$X = \text{Exp}_{\Sigma}(\text{Unvect}_{\Sigma}(t)), \quad t \sim \mathcal{N}(0, S).$$

The idea is to linearize the curved space $\mathcal{P}_d$ at the point Σ using the tangent space $T_{\Sigma} \mathcal{P}_d$ and to draw the Gaussian distribution $\mathcal{N}(0, S)$ on this flat space. Then, to respect the natural curvature of $\mathcal{P}_d$, the Exp_{Σ} wraps this flat distribution back onto the manifold. The parameter Σ is the center of the distribution and S is the dispersion matrix in the tangent space $T_{\Sigma} \mathcal{P}_d$. This wrapped Gaussian distribution has been studied in depth in [17] and several learning algorithms were derived in [18]. One can compute the density of this distribution in a closed form, which we omit here (see Theorem 4.2 of [17]). Given a sample $(\Sigma_1, \dots, \Sigma_N)$ of SPD matrices following a wrapped Gaussian $\text{WG}(\Sigma, S)$, one can estimate Σ using the Riemannian mean $\hat{\Sigma}_N$ (see Equation 2) and estimate S as follows:

$$\hat{S}_N = \frac{1}{N} \sum_{i=1}^N \text{VLog}_{\hat{\Sigma}_N}(\Sigma_i) \text{VLog}_{\hat{\Sigma}_N}(\Sigma_i)^{\top},$$

where $\text{VLog}_{\hat{\Sigma}_N}(\Sigma_i) = \text{Vect}_{\hat{\Sigma}_N}(\text{Log}_{\hat{\Sigma}_N}(\Sigma_i))$. These estimators are the maximum likelihood estimators.

One-Class Riemannian classifiers: We consider two types of one-class Riemannian classifiers. They are illustrated at Figure 1.

One-Class MDM: The first one-class classifier that we consider in this study is based on the *Minimum Distance to the Mean* (MDM) classifier introduced in [5]. We call this one-class classifier the *One-Class MDM* [6]. Given a training dataset $\Sigma_1, \dots, \Sigma_N$ of N covariance matrices of the class awake, the *One-Class MDM* computes their Riemannian mean $\hat{\Sigma}_N$ using equation 2. Then, given a new

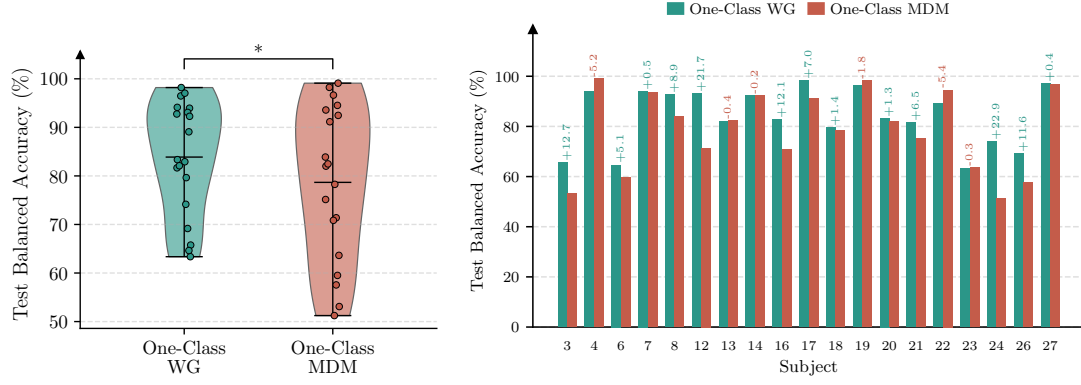
covariance matrix $\tilde{\Sigma}$, it computes the Riemannian distance $\delta_R(\tilde{\Sigma}, \hat{\Sigma}_N)$ between $\tilde{\Sigma}$ and $\hat{\Sigma}_N$ and classifies $\tilde{\Sigma}$ as awake if $\delta_R(\tilde{\Sigma}, \hat{\Sigma}_N) < \varepsilon$ and anesthetized otherwise, where ε is a threshold. In the experiments, the threshold ε was set to the median of the distances of the training data to $\hat{\Sigma}_N$ plus three times the standard deviation (same threshold as the Riemannian potato outlier detection method [19]).

One-Class WG (proposed): Now we introduce a novel one-class classifier, called *One-Class WG*, based on the Wrapped Gaussian distribution introduced in section *Wrapped Gaussian on $\mathcal{P}_d$* , which does not rely solely on the Riemannian mean of the training data. Given training covariance matrices $\Sigma_1, \dots, \Sigma_N$ of N from the awake class, our classifier estimates the parameters Σ and S of the wrapped Gaussian $\text{WG}(\Sigma, S)$ using the estimators $\hat{\Sigma}_N$ and $\hat{S}_N$. For a new covariance matrix $\tilde{\Sigma}$, the log-likelihood of $\tilde{\Sigma}$ under $\text{WG}(\hat{\Sigma}_N, \hat{S}_N)$ is computed and the trial is classified as awake if it exceeds a given threshold ε and anesthetized otherwise. In the experiments, the threshold ε was set to 95% of the training log-likelihoods. Notably, if the dispersion matrix is constrained to $S = \alpha I$, the classification becomes equivalent to a distance-based classifier. Thus, *One-Class MDM* is a particular isotropic case of the proposed *One-Class WG* model [18].

RESULTS

Main results: Let us start by comparing the mean balanced accuracy across all subjects of both classifiers *One-Class MDM* and *One-Class WG* on the data described in section *Materials and Methods*. The choice of this metric was driven by the highly asymmetric class imbalance. *One-Class MDM* achieved $78.67\% \pm 15.93$, whereas *One-Class WG* reached $83.88\% \pm 11.69$ (see Fig. 3a). A two-sided paired t-test confirmed that *One-Class WG* significantly outperformed *One-Class MDM* ($p = 0.0118$). Fig. 3b shows the accuracy of both classifiers for each subject, and for 13 subjects out of the 19, *One-Class WG* substantially improves the performances. These results were obtained when considering only five electrodes (Fig. 2) and for the *One-Class WG*, the estimated dispersion matrix $\hat{S}_N$ is only diagonal. We conducted an ablation study on the number of electrodes, on whether $\hat{S}_N$ should be diagonal, and on the choice of the log-likelihood threshold. Although performance varied only moderately across $\varepsilon \in \{0.85, 0.90, 0.95, 0.99\}$, significant differences between *One-Class WG* and *One-Class MDM* were obtained only at $\varepsilon = 0.95$.

Influence of the number of electrodes: We studied five different setups with an increasing number of electrodes (see Fig. 2). The results for the different setups, as well as the significance level of the two-sided paired t-test, with a Benjamini-Hochberg correction, between *One-Class WG* and *One-Class MDM* are given in Figure 4. One can see that, for the proposed *One-Class WG*, the best number of electrodes is 5, while for the *One-Class MDM*, the best number of electrodes is 11. With more electrodes, *One-*



(a) Distribution of the balanced accuracy across all subjects using 5 electrodes. A two-sided paired t-test shows that *One-Class WG* significantly outperforms *One-Class MDM* (* $p < 0.05$).

Figure 3: Balanced test accuracy of the two one-class classifiers (*One-Class WG* vs *One-Class MDM*) on the dataset. This is done on five electrodes (C3, C4, Cz, FC4 and FC3) and for the *One-Class WG*, the estimated Σ is only diagonal.

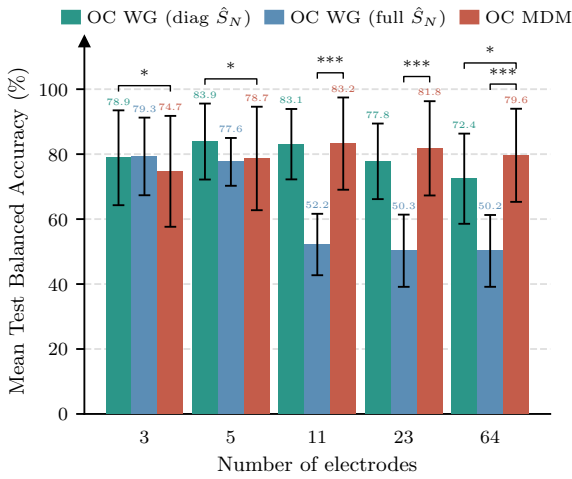


Figure 4: Mean test balanced accuracy of One-Class (OC) classifiers across different electrode setups. Statistical significance was assessed using a two-sided paired t-test (* $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$).

Class MDM outperforms our proposed method.

Influence of the type of the estimated $\hat{S}_N$: We also studied whether or not we should constrain the estimated dispersion matrix $\hat{S}_N$ to be diagonal: $\hat{S}_N = \text{diag}(\hat{s}_{N,1}^2, \dots, \hat{s}_{N,d(d+1)/2}^2)$. This parameterization reduces the number of parameters to estimate from $d(d+1)/2 \times d(d+1)/2$ to simply $d(d+1)/2$. The coefficients $\hat{s}_{N,j}^2$ are then simply the variances of the $(\text{Vect}_{\hat{S}_N}(\text{Log}_{\hat{S}_N}(\Sigma_i)))_i$. This leads to two variants of the proposed *One-Class WG*: with a full dispersion matrix $\hat{S}_N$ and with a diagonal one. The results of these two classifiers using the different montages are also in figure 4. One can see that, as soon as the number of electrodes is at least 5, the performances of the *One-Class WG* with a diagonal $\hat{S}_N$ are much better than the one with a full $\hat{S}_N$.

DISCUSSION

Overall, the proposed *One-Class WG* improved discrimination between awake and anesthetized conditions compared to the *One-Class MDM*, with reduced electrode montages, despite the variability of a real clinical setting. The following discussion examines these results in terms of continuous behavior of model scores across anesthetic depth, robustness to clinical variability, influence of montage, and the role of anisotropic dispersion modeling.

Continuous validation of one-class model scores through different levels of anesthetic depth:

To further understand model behaviors, we analyze the temporal evolution of their respective decision scores in relation to propofol concentration (see Fig. 5). Rather than focusing only on binary classification accuracy, this analysis allows us to evaluate how each model responds continuously to changes in anesthetic depth. *One-Class WG* score corresponds to the log-likelihood of each median nerve stimulation-evoked covariance matrix under the Wrapped Gaussian distribution. *One-Class MDM* score is the Riemannian distance to the awake centroid. Propofol concentration is displayed in the background to assess consistency with anesthetic depth. Both training (first 65 awake trials) and test data are displayed.

For subject 4 (94% *One-Class WG*, 99% *One-Class MDM*), both models clearly separate awake and anesthetized trials, with scores progressively diverging from the awake reference as propofol increases. The slightly lower *One-Class WG* accuracy may reflect a conservative decision threshold, rather than a failure of the underlying distributional modeling.

In contrast, subject 24 exhibits the largest discrepancy (74% *One-Class WG* vs. 51% *One-Class MDM*). While both models deviate from the awake reference beyond 4-5 $\mu\text{g/mL}$, *One-Class WG* yields a more consistent separation. One possible explanation lies in the EEG dynamics at high propofol concentrations, often associated with burst suppression patterns which are composed of periods of isoelectric activity and high-amplitude bursts [20]. Such dynamics induce strong variability in covari-

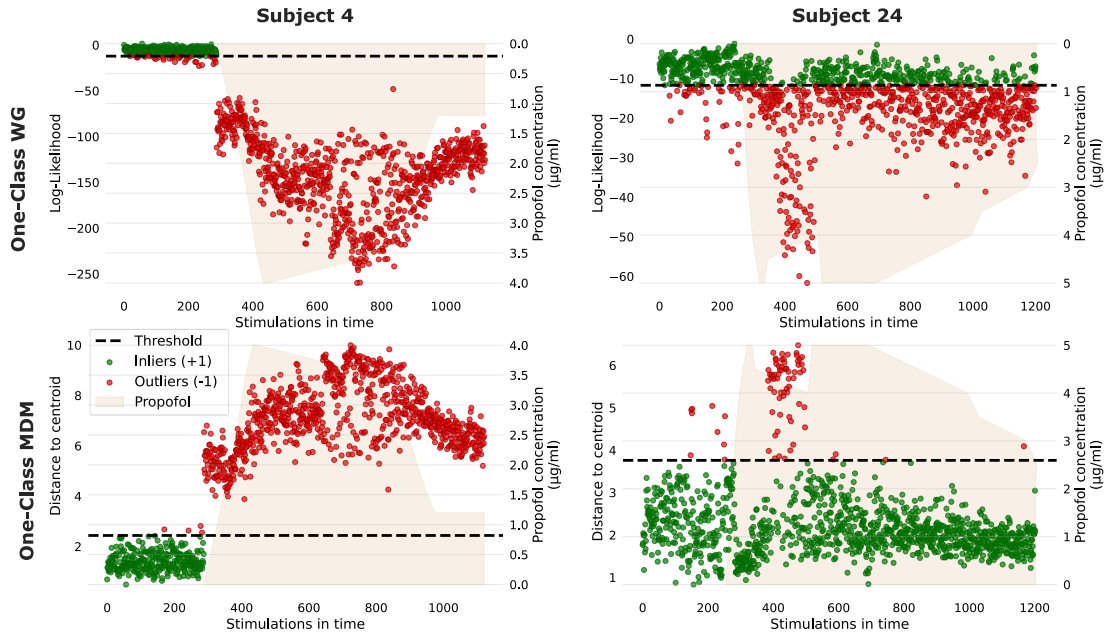


Figure 5: Evolution of one-class model scores over time, for two representative subjects: subject 4 (*One-Class WG*: 94%, *One-Class MDM*: 99%) and subject 24 (*One-Class WG*: 74%, *One-Class MDM*: 51%). These examples illustrate the inter-subject variability observed across the dataset.

ance structures. Because *One-Class MDM* relies solely on distance to a centroid, it does not explicitly model dispersion, whereas *One-Class WG* models not only the centroid but also the variability of the awake distribution through the dispersion matrix, yet covariance matrices corresponding to burst activity may not be uniformly distant from the awake centroid. Future analyses could quantify burst suppression via burst suppression ratio [12] and relate it to one-class model scores, or use Riemannian visualization techniques [21] to enable a geometric interpretation of how deep anesthesia shifts the data distribution relative to the awake reference. Further work will improve robustness in subjects exhibiting lower classification performance.

Robustness in the clinical anesthesia setting: An important aspect of the proposed approach is its robustness to variability. Unlike controlled laboratory BCI datasets, these data were acquired in a real surgical environment, introducing multiple sources of variability, including operating room noise, patient positioning affecting EEG stability, heterogeneous anesthesia protocols, and inter-subject differences. Despite this variability, *One-Class WG* outperformed *One-Class MDM*, suggesting that explicitly modeling dispersion of the reference class may improve robustness compared to centroid-based methods.

Effect of the montage: EEG was recorded with 64 channels, which leads to 64×64 covariance matrices. With *One-Class WG*, the estimated dispersion matrix $\hat{S}_N$ will be of dimension $2,080 \times 2,080$ (indeed, $64 * (64 + 1)/2 = 2,080$). The small number of training data (65 trials) and their short duration of 750 ms is not sufficient to correctly estimate $\hat{S}_N$ (see Section 5 of [17]). This is why we tried different montages with a lower number of electrodes. Moreover, in a real-life surgical setting, simpler setups are preferred. As shown in our results, the

proposed method works best with 3-5 electrodes, improving mean accuracy across all subjects by up to 5% over *One-Class MDM*. The best results are achieved with *One-Class WG* using only 5 electrodes and a diagonal $\hat{S}_N$.

Isotropy vs anisotropy: Let us now make an hypothesis to explain why the proposed *One-Class WG* performs better than *One-Class MDM*. *One-Class MDM* estimates only the mean of the training covariance matrices and therefore relies only on first order information. *One-Class WG*, estimates also a dispersion matrix, gathering second order information. The estimated dispersion matrix $\hat{S}_N$ captures the spread of the point cloud and may exhibit preferred directions, reflecting the anisotropic structure of the data rather than an isotropic one. This anisotropic behavior of the point cloud cannot be captured by the *One-Class MDM*, as it focuses only on first order information. To quantify this effect, we looked at the variance of the diagonal elements of $\hat{S}_N$ in the 5-electrodes setup and when $\hat{S}_N$ is only diagonal (where *One-Class WG* performs the best). Large variance across diagonal coefficients suggest an anisotropic point cloud, whereas similar diagonal coefficients indicate isotropy. Let us write the matrix $\hat{S}_N = \text{diag}(\hat{s}_{N,1}^2, \dots, \hat{s}_{N,d(d+1)/2}^2)$, then, we compute the variance of the diagonal coefficients ($\hat{S}_N$ has $d(d+1)/2$ diagonal coefficients), yielding one variance per subject. Subjects are then partitioned according to whether *One-Class WG* or *One-Class MDM* achieved higher accuracy, and we compute the mean variance within each group. Across all subjects, the group in which *One-Class WG* performs best has a mean variance of 3.021×10^{-2} , about 16 times larger than that of the group in which *One-Class MDM* performs best (1.891×10^{-3}). This suggests that *One-Class WG* tends to perform better when the point cloud is more anisotropic. Although this difference is partly driven by

a single outlier with unusually high variance, removing it still leaves the mean variance of the *One-Class WG* group 2.8 times larger. Statistical tests were not conducted due to the small number of subjects (6) in the latter group.

Perspectives: Several directions may extend this work. First, the clear separation between awake and anesthetized scores observed in Fig. 5 suggests that Riemannian change-point detection algorithms [22] could exploit the temporal structure of the data instead of treating trials independently. Second, the interpretability of Riemannian models could be assessed [23] as it remains a key aspect for clinical adoption. Future work will also include validation on additional EEG datasets, and extension to multimodal one-class models. Although no comparison with standard depth-of-anesthesia monitors was performed in the present study, such analyses are currently underway using a newly acquired dataset.

CONCLUSION

A new probabilistic one-class Riemannian classifier based on Wrapped Gaussians, *One-Class WG*, was introduced and evaluated in a BCI to detect anesthetic states. The proposed approach significantly outperformed the baseline *One-Class MDM* with a clinically-relevant number of EEG sensors, with a lower inter-subject variability, highlighting the importance of modeling the dispersion on the Riemannian manifold rather than relying solely on distance to a centroid. Notably, an optimal performance was achieved with fewer electrodes, making the approach particularly suitable for a practical BCI deployment in surgical environments. In summary, the proposed model offers a solution for one-class Riemannian BCI applications beyond the present anesthetic setting.

ACKNOWLEDGMENTS

This work was supported by the French National Research Agency (ANR), with projects PROTEUS (grant ANR-22-CE33-0015-01) and BCI4IA (grant ANR-22-CE19-0016-03).

REFERENCES

- [1] Lotte F et al. A review of classification algorithms for EEG-based brain-computer interfaces: A 10 year update. *J Neur Eng*. 2018.
- [2] Khan SS, Madden MG. One-class classification: Taxonomy of study and review of techniques. *CoRR*. 2013;abs/1312.0049.
- [3] Ruff L et al. A Unifying Review of Deep and Shallow Anomaly Detection. *Proc. IEEE*. 2021;109(5):756–795.
- [4] Yger F, Berar M, Lotte F. Riemannian Approaches in Brain-Computer Interfaces: A Review. *IEEE Trans Neur Syst Rehab*. 2017;25(10).
- [5] Barachant A, Bonnet S, Congedo M, Jutten C. Multiclass Brain-Computer Interface Classification by Riemannian Geometry. *IEEE Trans Biomed Eng*. 2012;59(4):920–928.
- [6] Marissens Cueva V et al. One-Class Riemannian EEG Classifier to Detect Anesthesia. In: *Proc. NEC*. 2024.
- [7] Said S, Bombrun L, Berthoumieu Y, Manton JH. Riemannian gaussian distributions on the space of symmetric positive definite matrices. *IEEE Trans. Inf. Theory*. 2017;63(4):2153–2170.
- [8] Rimbert S et al. Detection of motor cerebral activity after median nerve stimulation during general anesthesia (stim-motana): Protocol for a prospective interventional study. *JMIR Res. Protoc*. 2023;12:e43870.
- [9] Rimbert S, Riff P, Gayraud N, Schmartz D, Bougrain L. Median nerve stimulation based bci: A new approach to detect intraoperative awareness during general anesthesia. *Front. Neurosci*. 2019;13.
- [10] Schnitzler A, Salenius S, Salmelin R, Jousmäki V, Hari R. Involvement of primary motor cortex in motor imagery: A neuromagnetic study. *NeuroImage*. 1997;6(3):201–208.
- [11] Delorme A, Makeig S. Eeglab: An open source toolbox for analysis of single-trial eeg dynamics including independent component analysis. *J Neurosc Meth*. 2004;134(1):9–21.
- [12] Purdon PL, Sampson A, Pavone KJ, Brown EN. Clinical Electroencephalography for Anesthesiologists: Part I: Background and Basic Signatures. *Anesth*. 2015;123(4):937–960.
- [13] Gramfort A et al. Meg and eeg data analysis with mne-python. *Front Neuroinf*. 2013;7:267.
- [14] Pedregosa F et al. Scikit-learn: Machine learning in Python. *Journal of Machine Learning Research*. 2011;12:2825–2830.
- [15] Barachant A et al. Pyriemann. *Zenodo*. 2023.
- [16] Pennec X. Manifold-valued image processing with SPD matrices. In: *Riemannian Geometric Statistics in Medical Image Analysis*. Elsevier, 2020.
- [17] de Surrél T, Lotte F, Chevallier S, Yger F. Wrapped Gaussian on the manifold of symmetric positive definite matrices. In: *Proc. ICML*. 2025.
- [18] de Surrél T, Yger F, Lotte F, Chevallier S. A probabilistic view on riemannian machine learning models for SPD matrices. In: *Proc. GSI*. 2025.
- [19] Barthélemy Q, Mayaud L, Ojeda D, Congedo M. The Riemannian Potato Field: A Tool for Online Signal Quality Index of EEG. *IEEE Trans Neural Syst Rehab*. 2019;27(2):244–255.
- [20] Pawar N, Barreto Chang OL. Burst suppression during general anesthesia and postoperative outcomes: Mini review. *Front. in Syst. Neurosci.*. 2022;15.
- [21] de Surrél T, Chevallier S, Lotte F, Yger F. Geometry-aware visualization of high dimensional symmetric positive definite matrices. *TMLR*. 2025.
- [22] Wang X, Borsoi R, Breloy A, Richard C. Riemannian change point detection on manifolds with robust centroid estimation. 2025. eprint: 2508.18045 (cs.LG).
- [23] de Surrél T, Venot T, Corsi MC, Yger F. Interpretability of Riemannian tools used in brain computer interfaces. In: *IEEE MLSP*. 2025, 1–6.