

Inter-Expert Variability in Common Spatial Pattern Component Selection: Implication for Neurofeedback Applications

C. Dumas¹, C. Dussard¹, M.-C. Corsi^{1*}, N. George^{1*}

¹ Sorbonne Université, Institut du Cerveau – Paris Brain Institute – ICM, Inserm, CNRS, AP-HP, INRIA, CENIR MEG-EEG & NERV lab & Experimental Neurosurgery team, Hôpital de la Pitié Salpêtrière, Paris, France

* co-last authors

E-mail: Cassandra.dumas@icm-institute.org

ABSTRACT: Common Spatial Patterns (CSP) are widely used in motor imagery (MI)-based brain-computer interfaces (BCIs). While classification pipelines combine multiple components to optimize performance, MI-based neurofeedback (MI-NF) requires selecting a single spatial filter to drive feedback based on physiologically interpretable sensorimotor activity. This critical step is often implicit and relies on expert judgment. We quantified inter-expert variability in CSP selection using data from 20 subjects in a right-hand MI-NF experiment. Twenty-three BCI or neurophysiology experts independently selected the most physiologically relevant component among six CSP candidates per subject. Inter-expert agreement was fair overall (Fleiss' $\kappa = 0.256$) but varied substantially across subjects. Consensus did not systematically favor the top-ranked component, and selections spanned all six candidates. Although higher-ranked components elicited stronger agreement, substantial variability persisted. These findings identify single-component CSP selection as an overlooked source of methodological heterogeneity in MI-NF systems, highlighting the need for explicit and reproducible selection criteria.

INTRODUCTION

Common Spatial Pattern (CSP) decomposition is a widely used spatial filtering method in motor imagery (MI)-based brain-computer interfaces (BCIs) with electroencephalography (EEG) [1]. By maximizing variance differences between conditions, CSP enhances class separability. In classification-oriented BCIs, 2 or 3 pairs of components are typically selected to optimize discriminative performance, and multiple filters may be combined within machine learning pipelines [2]. In MI-based neurofeedback (MI-NF), however, the brain activity driving the feedback must reflect a physiologically interpretable neural process that subjects learn to voluntarily modulate [3]. The objective therefore extends beyond statistical discrimination to the extraction of a meaningful and controllable activity, typically expressed as event-related desynchronization (ERD) of the sensorimotor rhythm [4]. Traditional MI-NF commonly relies on predefined spatial filters, such as Laplacian derivations centered on C3/C4 electrodes, applied uniformly across subjects [3].

Although this ensures anatomical interpretability and methodological consistency, it overlooks substantial inter-individual variability in sensorimotor rhythm topography. CSP offers a subject-specific alternative. Yet, it typically yields multiple candidate components with high discriminative power. This creates the critical methodological challenge of determining which single component should drive feedback.

In practice, CSP selection may rely on quantitative criteria, including eigenvalues, mutual information, or cross-validation accuracy [5], or on qualitative neurophysiological judgment, such as the presence of focal sensorimotor activity consistent with motor cortex anatomy and known to be associated with MI [6]. However, identifying the most relevant component typically involves expert judgment, particularly when several components exhibit comparable discriminability or plausible spatial topography. This introduces a non-standardized, expert-dependent decision into system configuration.

Such expert-dependent selection may have broad implications. The selected spatial filter directly determines the neural signal fed back to the participant, thereby shaping the feedback representation and potentially influencing learning dynamics. Variability in expert selection translates into variability in feedback signals across studies, laboratories, and experimenters or clinicians, constituting an unrecognized source of methodological heterogeneity.

Previous research has examined numerous factors influencing MI-BCI performance, including signal processing choices, feedback modality, task instructions, and experimenter-participant interactions [8]. However, the influence of expert decisions on the internal configuration of NF systems remains largely unexplored. No study to date has systematically quantified inter-expert variability in CSP component selection for individual subjects in MI-NF.

The present study addresses this gap by characterizing agreement among 23 BCI and neurophysiology experts in selecting CSP components for right-hand MI-NF in 20 subjects. By quantifying inter-expert variability, this study empirically investigates the extent to which CSP selection reflects shared expert consensus versus expert-dependent methodological choice in face of inter-subject variability in MI-related ERD and CSPs.

MATERIALS AND METHODS

EEG Dataset

We analyzed the previously recorded EEG data from 20 healthy subjects (11 female; mean age $\pm$ SD: 25.8 ± 5.8 years). These data were acquired during a MI-NF CSP calibration protocol [9].

EEG was recorded with an actiCHamp Plus system (Brain Products GmbH) using a 32-active-electrode cap (ActiCAP snap) positioned according to the international 10–20 system. EEG signals were sampled at 1 kHz with a hardware band-pass filter of DC–280 Hz, referenced to Fz, with the ground at Fpz. Four electrodes (TP9, FT9, TP10, FT10) were excluded due to frequent artifact contamination, resulting in a final montage of 28 electrodes.

The calibration protocol included two conditions (rest and right-hand MI), each comprising 10 randomized trials of 8 s. During MI trials, subjects repeatedly imagined opening and closing their right hand while a virtual hand displayed similar movements on a computer screen, independent of neural activity.

CSP computation

To generate candidate spatial filters, EEG signals were notch-filtered at 50 Hz and band-pass filtered between 8 and 30 Hz, corresponding to the sensorimotor rhythm frequency band. No further artifact correction was applied, consistent with classical online BCI usage. CSP filters were estimated using the MNE-Python implementation [10]. Components were ranked according to their mutual information with class labels, and the six highest-rank CSP components were retained for each participant. For neurophysiological interpretation, CSP filters were converted to spatial patterns by projection into sensor space [2].

Expert Survey and CSP Selection Task

An anonymous online questionnaire was disseminated via learned societies, mailing lists, and specialized online communities. The call explicitly targeted researchers with expertise in either BCI or human electrophysiology; participation was voluntary and based on self-declared expertise. Twenty-three experts completed the survey. Non-BCI experts ($n=14$) were researchers with expertise in human EEG but without specific expertise in BCI systems. Among BCI experts, 6 out of 9 specifically worked with MI-based paradigms.

The questionnaire first collected background information, including field of expertise (open-ended), years of professional experience (< 5 years, 5–10 years, > 10 years), and self-rated familiarity with the CSP methodology (none, limited, moderate, high).

Experts then completed a CSP selection task for each of the 20 study subjects. The data for each subject were presented in a fixed randomized order identical across experts, and experts were blind to subjects' identifiers and characteristics. For each subject, experts viewed a composite figure including (i) a baseline-corrected ERD

topographical map computed from right-hand MI trials (power log-ratio relative to a -3 to -1 s baseline, averaged over 8–30 Hz and 0–8 s post-stimulus), and (ii) six candidate CSP components labeled A–F, corresponding to the six components ranked 1 to 6 by mutual information.

Experts were instructed to select the single CSP pattern they considered the most physiologically relevant for an online right-hand MI-NF application and to rate their confidence on a 5-point Likert scale (1 = very low confidence, 5 = very high confidence). Experts also indicated for each participant whether one CSP clearly stood out, several appeared equally plausible, or none appeared satisfactory. Instructions specified that CSP polarity is arbitrary and that no selection was considered correct or incorrect. Although CSP labels A–F reflected mutual-information ranking, the ranking criterion itself and its order were not explicitly disclosed, in order to encourage selections primarily based on neurophysiological interpretation rather than algorithmic metrics.

Inter-Expert Agreement and Association with CSP Rank

For each of the 20 subjects, we computed two metrics to quantify expert consensus across the six CSP candidates. First, inter-expert agreement was defined as the proportion of experts selecting the modal CSP (i.e., the most frequently selected component), ranging from 0 (no consensus) to 1 (perfect consensus). Second, choice dispersion was quantified using normalized Shannon entropy over the six CSP choices, ranging from 0 (perfect consensus) to 1 (maximal disagreement: uniform distribution across all six CSPs) [11].

Overall expert agreement across all 20 subjects was quantified using Fleiss' kappa (23 experts, 20 subjects, 6 CSPs), which measures inter-rater agreement beyond chance [12]. Kappa values were interpreted according to standard conventions: < 0 (no agreement), 0–0.20 (slight), 0.21–0.40 (fair), 0.41–0.60 (moderate), 0.61–0.80 (substantial), 0.81–1.00 (almost perfect). Ninety-five percent confidence intervals were estimated using non-parametric bootstrap resampling across subjects (10,000 iterations).

To examine whether consensus strength was related to the CSP rank, we computed, for each subject, the mean rank of the CSP selected by experts (rank = 1–6 according to mutual information). The association between mean selected rank and inter-expert agreement (proportion of modal CSP selection) was evaluated using Spearman correlation. We checked that agreement was not driven by CSP-unfamiliar experts by re-calculating the modal CSP after excluding the experts that reported no familiarity with CSP ($n = 7$).

Spatial Consistency of Selected CSP Patterns

To assess the spatial consistency of expert selections, the CSP components selected by each expert for each subject were normalized to unit Euclidean norm (L2 normalization) and sign-aligned so that the C3 electrode

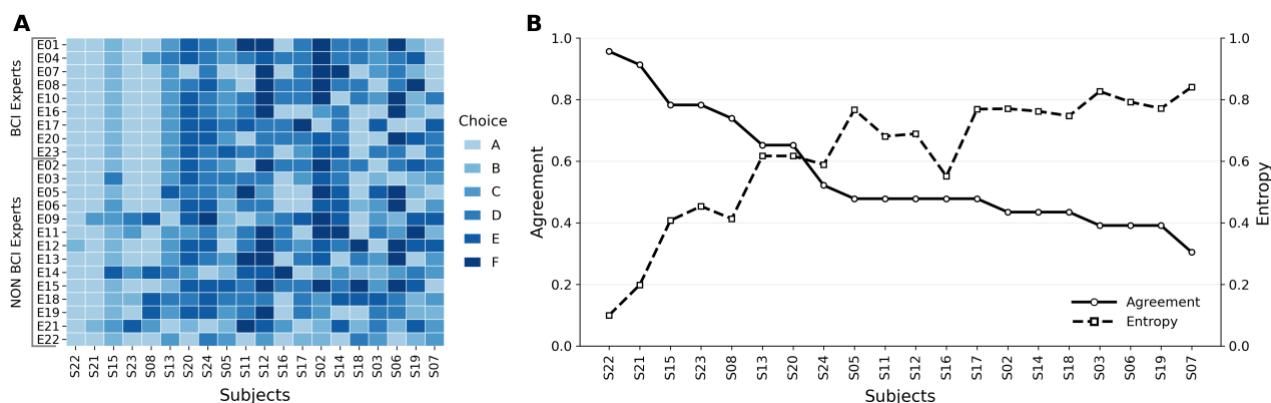


Figure 1: Inter-expert agreement on CSP selection. (A) Colormap showing the experts' (E01, E04, ... E22) component selections (A–F) for each subject (S22, S21, ... S07). Subjects are ordered from highest to lowest consensus (left to right columns). Experts are grouped according to their domain of expertise: BCI experts (top rows) and non-BCI experts (bottom rows). Colors represent the selected CSP component for each expert–subject pair. (B) Inter-expert agreement (in solid black line) and normalized Shannon entropy (in dashed black line) for each subject (in the same order as in (A)). The left y-axis represents agreement (modal selection proportion; 1 = full agreement, 0 = no agreement). The right y-axis represents entropy (0 = no dispersion, 1 = maximal dispersion across choices).

exhibited negative values, consistent with expected sensorimotor desynchronization [4]. For each subject, the mean and standard deviation (SD) of CSP values across experts, on each electrode, were computed to obtain topographical maps of inter-expert variability. Similarly, for each expert, the mean and SD of selected CSP values on each electrode across subjects were computed to assess within-expert spatial variability in CSP selection.

Influence of Expert Characteristics

At the expert level, we examined whether selection behavior was associated with research background, professional experience, and familiarity with the CSP methodology. For each expert, two outcome variables were computed: (i) individual agreement, defined as the proportion of selections matching the modal (consensus) CSP across the 20 subjects, and (ii) mean confidence, defined as the average confidence rating across all selections.

Experts were classified according to research background (BCI vs non-BCI) and group differences in expert agreement and mean confidence were evaluated using independent samples t-tests. Professional experience was categorized into two groups (<5 years vs ≥5 years) to ensure balanced group sizes, and group differences in agreement and confidence were assessed using independent samples t-tests. Familiarity with CSP was categorized into three ordered levels (none, limited, moderate–high), with moderate and high familiarity pooled together to ensure balanced group sizes. A linear regression analysis was conducted treating familiarity as an ordinal predictor to test for a linear trend in agreement and confidence across increasing familiarity levels.

Effect sizes were reported using Cohen's d for t-tests. Assumptions of normality (Shapiro–Wilk test) and homogeneity of variance (Levene's test) were verified prior to parametric testing and non-parametric tests were used as robustness checks where appropriate. The significance threshold was set at $\alpha = 0.05$. All analyses

were performed in Python 3.12. The analysis code is publicly available at: https://github.com/Dumas-C/CSP_Inter-Expert_Variability_Neurofeedback.

RESULTS

Inter-Expert Agreement and Association with CSP Rank

Across the 23 experts and 20 subjects, Fleiss' kappa indicated fair inter-rater agreement ($\kappa = 0.256$, 95% bootstrap CI [0.157–0.334]) according to standard benchmarks. At the subject level, inter-expert agreement (proportion of modal CSP selection) was on average 0.559 ± 0.187 , and normalized Shannon entropy was 0.618 ± 0.208 , indicating substantial variability in expert consensus strength across subjects (Fig. 1).

A significant negative association was observed between mean selected CSP rank and inter-expert agreement (Spearman $\rho = -0.48$, $p = 0.031$), indicating that stronger consensus tended to occur when experts converged on higher-rank CSP components.

Regarding decision clarity, experts indicated that one CSP clearly stood out for 6 subjects, several CSPs appeared equally plausible for 8 subjects, and no CSP appeared satisfactory for 6 subjects.

Spatial Consistency of Selected CSP Patterns

The subjects for whom inter-expert agreement was high were also those for whom SD across experts was low over the whole scalp, whereas subjects yielding lower inter-expert agreement were associated with larger variability (SD) across experts.

Importantly, this variability was not uniformly distributed but concentrated in frontal regions (Fig. 2A). At the subject level, inter-expert variability was approximately twofold higher over frontal electrodes (Fp1: 0.225 ± 0.100 ; Fp2: 0.190 ± 0.100) than over C3 (0.106 ± 0.043 ; mean $\pm$ SD across subjects).

At the expert level, selected CSP components were globally consistent across subjects, with mean spatial patterns predominantly centered over the left sensorimotor cortex. However, intra-expert variability across subjects was again particularly high over frontal electrodes (Fp1: 0.242 ± 0.082 ; Fp2: 0.228 ± 0.063 ; for comparison: C3: 0.155 ± 0.019 ; mean $\pm$ SD across experts), particularly among experts showing low agreement with the modal selection (Fig. 2B).

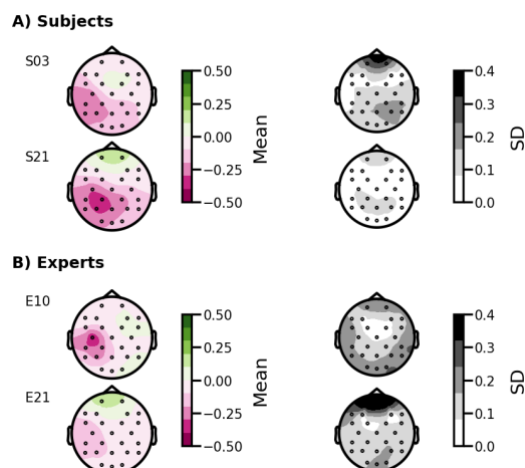


Figure 2: Topographical representation of the selected CSPs. The CSPs are displayed using a green–pink color scale, with green indicating positive weights and pink indicating negative weights. The variability (SD) of the selected CSPs is represented in grey maps. (A) Mean and SD of the CSPs selected by the experts for two representative subjects (S03 and S21). (B) Mean and SD of the CSPs selected by two representative experts (E10 and E21) across all subjects. This figure illustrates both inter-subject and inter-expert variability in the spatial distribution of the selected CSPs.

Influence of Expert Characteristics

Experts with BCI-related backgrounds ($n = 9$) demonstrated significantly greater individual agreement with the modal selection ($t(21) = 2.48$, $p = 0.022$, $d = 1.03$) and higher mean confidence ($t(21) = 2.87$, $p = 0.009$, $d = 1.20$) relative to non-BCI experts ($n = 14$). These effects were confirmed using non-parametric tests. Note, however, that the two groups differed in size and were small, which limits the statistical power of this comparison. Professional experience was categorized into two groups (<5 years, $n = 13$; ≥ 5 years, $n = 10$). No significant difference between these groups was observed for mean confidence ($t(21) = -0.88$, $p = 0.390$, $d = -0.38$). Individual agreement showed a non-significant trend toward lower values in less experienced experts ($t(21) = -2.01$, $p = 0.057$, $d = -0.86$). Familiarity with CSP technique was treated as an ordinal variable with three levels (none, $n = 7$; limited, $n = 9$; moderate/high, $n = 7$). Familiarity showed a significant positive linear trend with both individual agreement ($p = 0.037$) and confidence ($p = 0.037$) (Fig. 3).

When excluding CSP-unfamiliar experts ($n = 7$), the

modal CSP selection remained identical in 19 of 20 subjects (95%) ruling out that these CSP-unfamiliar experts could account for the results.

DISCUSSION

CSP is a widely used method in BCI applications, which enables defining individualized spatial features for feedback. While it may help refining the targeting of neural activity in MI-NF, the use of CSP in MI-NF requires selecting a single physiologically interpretable component [2]. This selection constitutes a critical methodological decision, as the chosen filter determines which neural process participants are trained to modulate. The present study quantifies inter-expert variability in CSP component selection to evaluate the level of expert consensus and its potential implications for MI-NF.

Fleiss' kappa indicated fair agreement ($\kappa = 0.256$), demonstrating that single-component CSP selection is neither trivial nor fully standardized. Importantly, in multi-category settings involving six response options and 23 independent raters, κ statistics are known to be conservative and sensitive to marginal distributions [13]. Thus, the observed value reflects substantial, structured agreement well beyond chance. At the subject level, expert agreement strength varied markedly: Some subjects elicited strong convergence, but most generated distributed selections across two or three plausible candidates, leading to a level of inter-expert agreement below .5 (less than half the experts selecting the modal component) in a majority of subjects. Normalized entropy analysis further showed that identical modal agreement proportions could mask distinct underlying selection structures. Together, these findings suggest that CSP selection operates as a semi-structured expert decision guided by partially shared evaluation principles rather than as a deterministic outcome of signal processing.

Inter-expert agreement showed a negative association with the mean selected CSP rank. This indicates that stronger consensus tended to emerge when experts converged on higher-rank components. This supports partial alignment between algorithmic discriminability and physiological plausibility. In subjects exhibiting clear contralateral sensorimotor modulation, the top-ranked CSP component appeared to coincide with physiologically meaningful activity [1]. However, this alignment was not systematic. In several cases, experts preferentially selected lower-ranked components, which had actually a reduced discriminative value between MI and rest in terms of mutual information, a measure often used for automatic selection of CSP in BCI settings. This behavior likely reflects situations in which top-rank components captured discriminative variance that was spatially diffuse or physiologically ambiguous. Rather than indicating a limitation of the CSP algorithm, these findings suggest that algorithmic ranking does not necessarily coincide with interpretability constraints

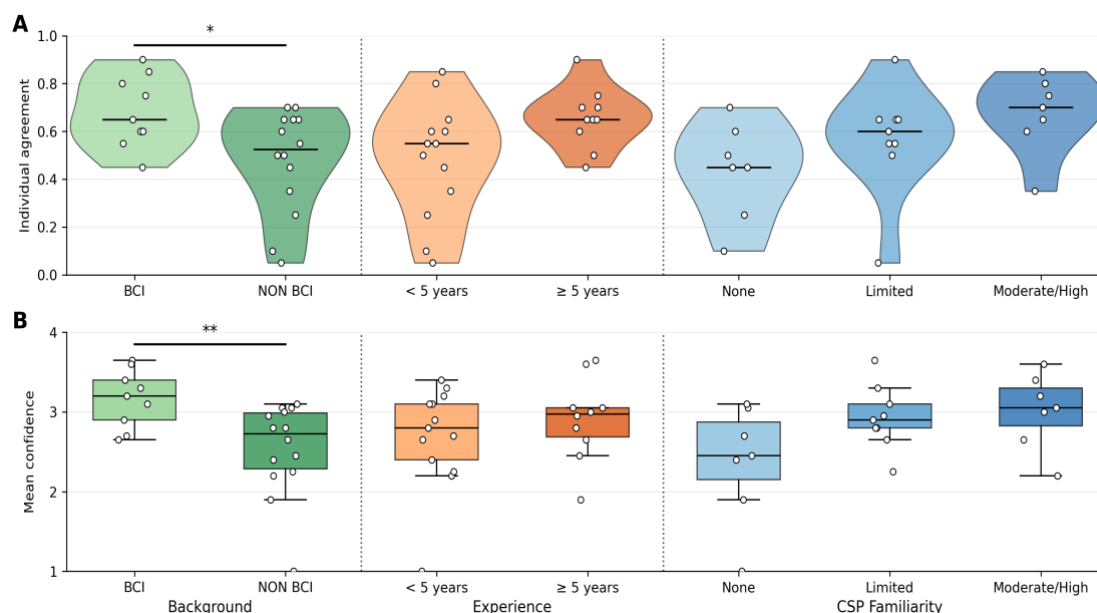


Figure 3: Influence of expert characteristics on their individual agreement and mean confidence. The individual expert agreement with the consensus (aka. modal CSP selection) and the mean confidence is compared across three expert characteristics: background (BCI vs. non-BCI, left plot in green), years of experience (< 5 vs. ≥ 5 years, middle plot in orange), and CSP familiarity (none, limited, moderate-to-high, right plot in blue). In each plot, white dots represent individual expert data points. (A) Individual expert agreement, defined as the proportion of subjects for whom an expert’s selected the consensus (aka. modal) CSP component. Violin plots show the distribution within each expert group; the central black line represents the median. (B) Mean confidence ratings of the experts. Boxplots display the interquartile range (Q1–Q3), the median (central black line), and whiskers extend to 1.5× IQR (with outliers excluded).

specific to NF applications.

Spatial consistency analysis further clarified sources of disagreement: Variability was predominantly concentrated over frontal regions. These regions are notably associated with ocular artifacts CSP decomposition is known to be sensitive to variance-dominant sources, including artifacts [14]. Although the present study did not directly validate artifact contributions, the observed spatial distribution is consistent with the hypothesis that ambiguous or artifact-prone components contribute to interpretative divergence. This raises practical considerations for real-time NF implementation. Incorporating online artifact correction strategies or physiologically constrained spatial filtering approaches may enhance the specificity of extracted components and increase inter-expert consensus [15]. Future studies should evaluate preprocessing effects on CSP selection variability.

Experts with BCI-focused research backgrounds reported higher confidence and demonstrated greater alignment with modal selections. This difference may reflect greater familiarity with MI datasets and CSP-derived spatial representations rather than superior neurophysiological knowledge. Treating familiarity as an ordered variable revealed a significant positive linear trend for both individual agreement and confidence. This indicates that greater familiarity with CSP is associated with more consistent and confident selection behavior. However, even among the experts familiar with CSPs, substantial variability persisted. When excluding CSP-unfamiliar experts, modal selections remained identical

in 19 of 20 subjects, confirming that observed variability was not driven by methodological inexperience. Moreover, experts reported that either several CSPs appeared equally plausible or none appeared satisfactory for 14 of 20 subjects, indicating that selection frequently involved navigating ambiguous candidates rather than identifying a single unequivocal solution. These findings have implications for methodological transparency and reproducibility in MI-NF studies. In practice, different laboratories may implement similar systems yet select distinct spatial filters, thereby training subjects on different neural signals. Whether such configuration differences influence learning and/or behavioral outcomes remains an open empirical question. It is noteworthy that, although CSP is widely used in MI-based classification BCIs, its use as a single spatial filter to directly drive neurofeedback remains relatively uncommon [16]. Most MI-NF paradigms rely on predefined anatomical filters. The need to select one physiologically interpretable component among several discriminative candidates may partly explain this limited adoption. Nonetheless, in a domain already characterized by substantial inter-individual variability and NF/BCI inefficiency [7], unreported variability in expert-driven system configuration may contribute to inconsistent outcomes across studies. Systematic reporting of CSP ranks, selection rationale, spatial topographies, time interval between calibration and real time application, and, when possible, open sharing of spatial filters would represent a minimal yet important step toward improving transparency and cross-study comparability.

A limitation of the present study is that it characterizes variability in expert choices without directly evaluating the functional consequences of these choices on NF performance. While we demonstrate substantial inter-expert variability in spatial filter selection, we do not assess if different CSP selections would lead to differences in subjects' learning or behavioral outcomes. It remains possible that multiple candidate components, despite spatial differences, yield comparable NF efficacy. Conversely, it is also plausible that certain selections enhance controllability or engagement of task-relevant neural processes. Determining whether inter-expert variability translates into meaningful differences in NF performance constitutes a critical next step.

More broadly, the observed pattern of structured but incomplete agreement suggests that experts rely on partially shared interpretative criteria when evaluating CSP components. Formalizing these criteria into explicit, reproducible guidelines or developing hybrid algorithmic approaches that integrate discriminability with physiologically informed constraints may help reduce latent configuration variability in MI-based NF / BCI systems. In particular, data-driven approaches such as transfer learning could reduce reliance on manual expert selection by leveraging across-subject reliable priors to identify physiologically relevant components more robustly. Such efforts could strengthen interpretability, reproducibility, and translational robustness of CSP-based NF approaches.

CONCLUSION

This study demonstrates that single-component CSP selection in MI-NF is not fully standardized and yields only fair inter-expert agreement. Although consensus often emerges for higher-rank components, substantial variability persists, particularly in spatially ambiguous cases. These findings identify CSP selection as a source of configuration-level variability in MI-NF systems. Implementing on-line physiological (ocular) artifact correction may help reduce this variability. Formalizing or automating the CSP selection step, and reporting selected components, would contribute to transparency and reproducibility in NF and BCI research.

ACKNOWLEDGEMENTS

This work was supported by the French National Research Agency (BETAPARK project, ANR-20-CE37-0012). The authors thank the BCI and neurophysiology experts who participated in the survey for their valuable contribution to this study.

REFERENCES

[1] Ramoser H, Müller-Gerking J, Pfurtscheller G. Optimal spatial filtering of single trial EEG during imagined hand movement. *IEEE Trans Rehabil Eng.* 2000;8(4): 441-446
[2] Blankertz B, Tomioka R, Lemm S, Kawanabe M,

Müller KR. Optimizing spatial filters for robust EEG single-trial analysis. *IEEE Signal Process Mag.* 2008;25(1): 41-56
[3] Marzbani H, Marateb HR, Mansourian M. Neurofeedback: a comprehensive review on system design, methodology and clinical applications. *Basic Clin Neurosci.* 2016;7(2): 143-158
[4] Pfurtscheller G, Lopes da Silva FH. Event-related EEG/MEG synchronization and desynchronization: basic principles. *Clin Neurophysiol.* 1999;110(11): 1842-1857
[5] Ang KK, Chin ZY, Wang C, Guan C, Zhang H. Filter bank common spatial pattern algorithm on BCI competition IV datasets 2a and 2b. *Front Neurosci.* 2012;6: 39
[6] Bennett JD, John SE, Grayden DB, Burkitt AN. A neurophysiological approach to spatial filter selection for adaptive brain-computer interfaces. *J Neural Eng.* 2021;18(2): 026006
[7] Zhang R, Li F, Zhang T, Yao D, Xu P. Subject inefficiency phenomenon of motor imagery brain-computer interface: influence factors and potential solutions. *Brain Sci Adv.* 2021;6(3): 224-241
[8] Pillette L, Roc A, N'Kaoua B, Lotte F. Experimenters' influence on mental-imagery based brain-computer interface user training. *Int J Hum Comput Stud.* 2021;149: 102603
[9] Dussard C, Pillette L, Dumas C, Pierrieau E, Hugueville L, Lau B, et al. Influence of feedback transparency on motor imagery neurofeedback performance: the contribution of agency. *J Neural Eng.* 2024;21(5): 056029
[10] Gramfort A, Luessi M, Larson E, Engemann DA, Strohmeier D, Brodbeck C, et al. MEG and EEG data analysis with MNE-Python. *Front Neuroinform.* 2013;7: 267
[11] Shannon CE. A mathematical theory of communication. *Bell Syst Tech J.* 1948;27(3): 379-423
[12] Fleiss JL. Measuring nominal scale agreement among many raters. *Psychol Bull.* 1971;76(5): 378-382
[13] Almehrizi RS. Coefficient Lambda for Interrater Agreement Among Multiple Raters: Correction for Category Prevalence. *Educ Psychol Meas.* 2025;00131644251380540
[14] Yong X, Ward RK, Birch GE. Robust common spatial patterns for EEG signal preprocessing. In: 2008 30th Annual International Conference of the IEEE Engineering in Medicine and Biology Society. Vancouver (BC), Canada; 2008. p. 2087-2090
[15] Dumas C, Corsi MC, Dussard C, Grosselin F, George N. Automatic ocular artifact correction in electroencephalography for neurofeedback. In: Proceedings of the 18th International Conference on Bio-inspired Systems and Signal Processing; 2025 Mar; Porto, Portugal
[16] Lioi G, Butet S, Fleury M, Bannier E, Lécuyer A, Bonan I, et al. A multi-target motor imagery training using bimodal EEG-fMRI neurofeedback: A pilot study in chronic stroke patients. *Front Hum Neurosci.* 2020;14:37