

Spatial blind source separation

BY FRANÇOIS BACHOC

*Institut de Mathématiques de Toulouse, Université Paul Sabatier,
118 route de Narbonne, 31062 Toulouse, France
francois.bachoc@math.univ-toulouse.fr*

MARC G. GENTON

*Statistics Program, King Abdullah University of Science and Technology,
Thuwal 23955-6900, Saudi Arabia
marc.genton@kaust.edu.sa*

KLAUS NORDHAUSEN

*CSTAT, Vienna University of Technology,
Wiedner Hauptstr. 7, A-1040 Vienna, Austria
klaus.nordhausen@tuwien.ac.at*

ANNE RUIZ-GAZEN

*Toulouse School of Economics, University of Toulouse Capitole,
21 allée de Brienne, 31000 Toulouse, France
anne.ruiz-gazen@tse-fr.eu*

AND JONI VIRTÄ

*Department of Mathematics and Statistics, University of Turku,
20014 Turun yliopisto, Finland,
Department of Mathematics and Systems Analysis, Aalto University,
PL 11000, 00076 AALTO, Finland.
joni.virta@utu.fi*

SUMMARY

Recently a blind source separation model was suggested for spatial data together with an estimator based on the simultaneous diagonalisation of two scatter matrices. The asymptotic properties of this estimator are derived here and a new estimator, based on the joint diagonalisation of more than two scatter matrices, is proposed. The asymptotic properties and merits of the novel estimator are verified in simulation studies. A real data example illustrates the method.

Some key words: Joint diagonalisation; Limiting distribution; Multivariate random field; Spatial scatter matrix.

1. INTRODUCTION

There is an abundance of multivariate data measured at spatial locations $s_1, \dots, s_n$ in a domain $\mathcal{S}^d \subseteq \mathbb{R}^d$. Such data exhibit two kinds of dependence: measurements taken closer to each other

tend to be more similar than measurements taken further apart, and the variable values within a single location are likely to be correlated.

This complexity makes modelling multivariate spatial data computationally and theoretically difficult due to the large number of parameters required to represent the dependencies. In this work we address this problem through blind source separation, a framework established as independent component analysis for independent and identically distributed data and for stationary and non-stationary time series; see Comon & Jutten (2010) and Nordhausen & Oja (2018). Denoting a p -variate random field as $X(s) = \{X_1(s), \dots, X_p(s)\}^T$, where T is the transpose operator, we assume that $X(s)$ obeys the spatial blind source separation model introduced in Nordhausen et al. (2015). That is, $X(s)$ at a location s is a linear mixture of an underlying p -variate latent field $Z(s) = \{Z_1(s), \dots, Z_p(s)\}^T$ with independent components,

$$X(s) = \Omega Z(s), \quad (1)$$

where Ω is an unknown $p \times p$ full rank matrix. In this introduction section, we consider that the random fields X and Z have mean functions zero, for the sake of simplicity.

When the observed random field X takes the form (1), modeling and computational simplifications can be obtained. Indeed, if no assumption at all is made on X , then the distribution of X is characterized by p covariance functions and by $p(p-1)/2$ cross-covariance functions. In contrast, when it is assumed that X takes the form (1), then the distribution of X is characterized by p covariance functions and by a $p \times p$ matrix. As a function is an infinite-dimensional object, it is more difficult to model and estimate than a fixed-dimensional matrix. Thus, when the observed random field X takes the form (1), modeling simplifications are available.

When no assumption is made on X , a common practice in geostatistics is to let each of the p covariance functions and each of the $p(p-1)/2$ cross-covariance functions of X be characterized by q parameters. For instance the case $q = 2$ can correspond to a variance and a length scale parameter for an isotropic function. Then, the resulting $qp(p+1)/2$ parameters are usually estimated jointly by optimizing a fit criterion, typically the likelihood (Genton & Kleiber, 2015). This requires to perform an optimization in dimension $qp(p+1)/2$, where the computational cost of an evaluation of the likelihood is $O(p^3n^3)$. Once the $qp(p+1)/2$ parameters are estimated, the prediction of $X(s)$ for new values of s can be performed at the computational cost $O(p^3n^3)$.

In contrast, consider that model (1) holds for X . We will show in this paper that an estimate of Ω^{-1} can be obtained. This is carried out by, first, computing scatter matrices with computational cost $O(p^2n^2)$ and, second, performing an optimization in dimension p^2 where the computational cost of the function to be evaluated is $O(p^2)$, see § 4 for details. If each covariance function of Z is characterized by q parameters, each of them can be estimated separately, by optimizing the likelihood in dimension q . The evaluation cost of the likelihood is $O(n^3)$. Once the qp covariance parameters are estimated, the prediction of $X(s)$ for new values of s can be performed at cost $O(pn^3)$. Indeed, the predictions of $Z_1(s), \dots, Z_p(s)$ can be performed separately at cost $O(n^3)$ and aggregated with negligible cost.

Not all random fields X obey a spatial blind source separation model of the form (1). For instance, (1) forces the cross-covariance functions of X to be symmetric. Nevertheless, it is a reasonable assumption in a fair number of practical situations (Nordhausen et al., 2015) and brings the computational benefits discussed above. Furthermore, an additional benefit of the form (1) is dimension reduction. In blind source separation, often significantly fewer than the full p latent components are needed to capture the essential structure of the original observations and the remaining components can be discarded as noise.

We thus consider the spatial blind source separation model (1) in this paper and focus on the estimation of Ω^{-1} . As discussed above, this estimation enables to estimate the cross-covariance functions of X and to perform prediction. Our approach for estimating Ω^{-1} is based on the use of local covariance, or scatter, matrices,

$$\widehat{M}(f) = n^{-1} \sum_{i=1}^n \sum_{j=1}^n f(s_i - s_j) X(s_i) X(s_j)^T, \quad (2)$$

where $f : \mathbb{R}^d \rightarrow \mathbb{R}$ is called the kernel function. Nordhausen et al. (2015) obtained estimators $\widehat{\Gamma}(f)$ of Ω^{-1} through a generalized eigendecomposition of pairs of local covariance matrices with kernels of the form (f_0, f_h) , with $f_h(s_i - s_j) = I(\|s_i - s_j\| \leq h)$, for a positive constant h , where $I(\cdot)$ denotes the indicator function and $f_0(s) = I(s = 0)$. Their estimators were based on the following definition, with $f = f_h$ for some $h > 0$.

DEFINITION 1. *An unmixing matrix estimator $\widehat{\Gamma}(f)$ jointly diagonalizes $\widehat{M}(f_0)$ and $\widehat{M}(f)$ in the following way*

$$\widehat{\Gamma}(f) \widehat{M}(f_0) \widehat{\Gamma}(f)^T = I_p \quad \text{and} \quad \widehat{\Gamma}(f) \widehat{M}(f) \widehat{\Gamma}(f)^T = \widehat{\Lambda}(f),$$

where $\widehat{\Lambda}(f)$ is a diagonal matrix with diagonal elements in decreasing order.

This method is conceptually close to principal component analysis where latent variables that have maximal variance are found through the diagonalisation of the covariance matrix. However, since the covariance matrix does not capture spatial information, it was extended to the concept of a local covariance matrix in Nordhausen et al. (2015). Analogously, diagonalising local covariance matrices then aims to find latent fields that maximize spatial correlation.

Here, we expand on their work by not restricting the kernel f in Definition 1 to be of the “ball” form f_h . Furthermore, we derive the asymptotic behavior for the method proposed in Nordhausen et al. (2015) for a large class of kernel functions f .

The idea when constructing these kernel functions is that the mean values of $\widehat{M}(f)$ and $\widehat{M}(f_0)$ would be diagonal matrices if, in their definition, the mixed components X were replaced by the latent components Z . Hence, a general blind source separation strategy is to undo the mixing in X by finding a matrix $\widehat{\Gamma}(f)$ which simultaneously diagonalizes $\widehat{M}(f)$ and $\widehat{M}(f_0)$. This is computationally simple and can always be done exactly using generalized eigenvalue-eigenvector theory. From temporal blind source separation, it is however well known that when diagonalising only two matrices, the choice of the matrices can have a large impact on the separation efficiency. Therefore, it is a popular strategy to approximately diagonalize more than two matrices with the hope of including more information; see for example Belouchrani et al. (1997), Nordhausen (2014), Miettinen et al. (2014), Matilainen et al. (2015) and Miettinen et al. (2016). Approximate diagonalization becomes then necessary as the matrices commute only at the population level but not when estimated using finite data. There are many algorithms available for this purpose. We use this idea to extend the method of Nordhausen et al. (2015) to jointly diagonalize more than two local covariance matrices. We also derive the asymptotic behaviour of these novel estimators.

2. SPATIAL BLIND SOURCE SEPARATION MODEL

2.1. General assumptions

In the spatial blind source separation model, the following assumptions are made:

Assumption 1. $E\{Z(s)\} = 0$ for $s \in \mathcal{S}^d$,

Assumption 2. $\text{cov}\{Z(s)\} = \mathbb{E}\{Z(s)Z(s)^\top\} = I_p$;

Assumption 3. $\text{cov}\{Z(s_1), Z(s_2)\} = \mathbb{E}\{Z(s_1)Z(s_2)^\top\} = D(s_1, s_2)$, where D is a diagonal matrix whose diagonal elements depend only on $s_1 - s_2$.

Let $\text{cov}\{Z_k(s_i), Z_k(s_j)\} = K_k(s_i - s_j) = D(s_i, s_j)_{k,k}$, where K_k denotes the stationary covariance function of Z_k , for $k = 1, \dots, p$.

Assumption 1 is made for convenience and can easily be replaced by assuming a constant unknown mean (see Lemma B.8 in the supplementary material). Assumption 2 says that the components of $Z(s)$ are uncorrelated and implies that the variances of the components are one, which reduces identifiability issues and comes without loss of generality. Assumption 3 says that there is also no spatial cross-dependence between the components. However, even after these assumptions are made, the model is not uniquely defined. The order of the latent fields and also their signs can be changed. This is common for all blind source separation approaches and is not considered a problem in practice.

2.2. Identifiability

The expectations of $\widehat{M}(f)$ and $\widehat{M}(f_0)$ are respectively

$$M(f) = n^{-1} \sum_{i=1}^n \sum_{j=1}^n f(s_i - s_j) \mathbb{E}\{X(s_i)X(s_j)^\top\} \quad \text{and} \quad M(f_0) = n^{-1} \sum_{i=1}^n \mathbb{E}\{X(s_i)X(s_i)^\top\}. \quad (3)$$

Thus the empirical procedure of Definition 1, operating on $\widehat{M}(f)$ and $\widehat{M}(f_0)$, can be associated to the following theoretical procedure, operating on $M(f)$ and $M(f_0)$.

DEFINITION 2. *For any function $f : \mathbb{R}^d \rightarrow \mathbb{R}$, an unmixing matrix functional $\Gamma(f)$ is defined as a functional which jointly diagonalizes $M(f)$ and $M(f_0)$ in the following way*

$$\Gamma(f)M(f_0)\Gamma(f)^\top = I_p \quad \text{and} \quad \Gamma(f)M(f)\Gamma(f)^\top = \Lambda(f),$$

where $\Lambda(f)$ is a diagonal matrix with diagonal elements in decreasing order.

We remark that an unmixing matrix $\Gamma(f)$ can be found using the generalized eigenvalue-eigenvector theory. In addition, an unmixing matrix is never unique, since if $\Gamma(f)$ and $\Lambda(f)$ satisfy Definition 2, then $S\Gamma(f)$ and $\Lambda(f)$ also satisfy Definition 2 for any diagonal matrix S with diagonal elements equal to -1 or 1 . We also remark that $\Lambda(f)$ is not the expectation of $\widehat{\Lambda}(f)$, in general. Indeed, Definitions 1 and 2 are based on non-linear functions of $\{\widehat{M}(f), \widehat{M}(f_0)\}$ and of $\{M(f), M(f_0)\}$.

The usual notion of identifiability in blind source separation is that any unmixing functional $\Gamma(f)$ should recover the components of Z up to signs and order of the components. Thus, any unmixing functional $\Gamma(f)$ should coincide with Ω^{-1} , up to the order and signs of the rows.

DEFINITION 3. *We say that the unmixing problem given by f is identifiable if any unmixing functional $\Gamma(f)$ satisfying Definition 2 can be written as $PS\Omega^{-1}$, where P is a permutation matrix and S is a diagonal matrix with diagonal elements equal to -1 or 1 .*

The motivation behind identifiability is that, if identifiability holds, then estimating $M(f_0)$ and $M(f)$ consistently by $\widehat{M}(f_0)$ and $\widehat{M}(f)$ enables to obtain $\widehat{\Gamma}(f)$, which will be approximately equal to a matrix of the form $PS\Omega^{-1}$, with P and S as in Definition 3. The following proposition provides a necessary and sufficient condition for identifiability. This proposition is proved in § B.2 of the supplementary material. All the other theoretical results in this paper are also proved in the supplementary material. Let $M^{-\top}$ denote the inverse of the transpose of M .

PROPOSITION 1. *The unmixing problem given by f is identifiable if and only if the diagonal elements of $\Omega^{-1}M(f)\Omega^{-T}$ are distinct.*

We remark that identifiability is a joint property of the kernel f and the covariance functions $K_1, \dots, K_p$. For instance, consider the situation where $K_1, \dots, K_p$ are compactly supported and equal to zero at distances larger than $0 < r < \infty$, and where one uses the function $f(s) = I(r_1 < \|s\| \leq r_2)$, with $r \leq r_1 < r_2 < \infty$ as kernel. Then identifiability does not hold because $\Omega^{-1}M(f)\Omega^{-T}$ is equal to the zero matrix. On the other hand, if f is a ball kernel of the form $f(s) = I(\|s\| \leq r_0)$ with $r_0 > 0$, then identifiability may hold, for the same covariance functions $K_1, \dots, K_p$.

Finally, for any kernel f , a necessary condition for identifiability is that there does not exist $k, l \in \{1, \dots, p\}$, $k \neq l$, such that $K_k(s_i - s_j) = K_l(s_i - s_j)$ for all $i, j = 1, \dots, n$. Indeed, if this was the case, then the diagonal elements k and l of $\Omega^{-1}M(f)\Omega^{-T}$ would be equal, for any kernel f . An extreme example of this issue is $K_1 = \dots = K_p$ with only Gaussian components. If this is the case, then, for any orthogonal matrix Q , the distribution of the random field QZ is the same as that of the random field Z . Hence, no statistical procedure can be expected to recover the components of Z , even up to signs and permutations, when only observing the transformed random field X .

2.3. Relationships with other models of multivariate random fields

The spatial blind source separation is notably different from the usual multivariate models for spatial data, which are often defined starting with their covariance functions contained in a cross-covariance matrix,

$$C(s_1, s_2) = \text{cov}\{X(s_1), X(s_2)\} := \{C_{k,l}(s_1, s_2)\}_{k,l=1}^p,$$

whereas our approach for estimating Ω^{-1} does not need to model or estimate the covariance functions of the latent fields $Z_1(s), \dots, Z_p(s)$.

In a recent extensive review, Genton & Kleiber (2015) discussed different approaches to define cross-covariance matrix functionals and gave a list of properties and conventions that they should satisfy, for instance stationarity and invariance under rotation. As Genton & Kleiber (2015) pointed out, to create general classes of models with well-defined cross-covariance functionals is a major challenge. Multivariate spatial models are particularly challenging as many parameters need to be fitted. In textbooks such as Wackernagel (2003) usually the following two popular models are described.

In the intrinsic correlation model it is assumed that the stationary covariance matrix $C(h)$ can be written as the product of the variable covariances and the spatial correlations, $C(h) = \rho(h)T$, for all lags h , where T is a non-negative definite $p \times p$ matrix and $\rho(h)$ a univariate spatial correlation function.

The more popular linear model of coregionalization is a generalization of the intrinsic correlation model, and the covariance matrix then has the form

$$C(h) = \sum_{k=1}^r \rho_k(h)T_k,$$

for some positive integer $r \leq p$ with all the ρ_k 's being univariate spatial correlation functions and T_k 's being non-negative definite $p \times p$ matrices, often called coregionalization matrices. Hence, with $r = 1$ this reduces to the intrinsic correlation model. The linear model of coregionalization implies a symmetric cross-covariance matrix.

Estimation in the linear model of coregionalization is discussed in several papers. Goulard & Voltz (1992) focused on the coregionalization matrices using an iterative algorithm where the spatial correlation functions are assumed to be known. The algorithm was extended in Emery (2010). Assuming Gaussian random fields, an expectation-maximisation algorithm was suggested in Zhang (2007) and a Bayesian approach was considered in Gelfand et al. (2004).

There is a simple connection between the spatial blind source separation model and the linear model of coregionalization. The covariance matrix $C_X(h)$ resulting from a spatial blind source separation model is always symmetric and can be written as

$$C_X(h) = \sum_{k=1}^p K_k(h) T_k,$$

with $T_k = \omega_k \omega_k^T$, ω_k being the k th column of Ω . Thus the spatial blind source separation model is a special case of the linear model of coregionalization with $r = p$ and where all coregionalization matrices T_k , $k = 1, \dots, p$, are rank one matrices.

3. ASYMPTOTIC PROPERTIES FOR SIMULTANEOUS DIAGONALISATION OF TWO MATRICES

Recall the definition (2) of a local covariance matrix and that

$$\widehat{M}(f_0) = n^{-1} \sum_{i=1}^n X(s_i) X(s_i)^T \quad (4)$$

is the covariance estimator. Asymptotic results can be derived for the previous estimators assuming that Assumptions 1 to 3 hold together with the following assumptions:

Assumption 4. The coordinates $Z_1, \dots, Z_p$ of Z are stationary Gaussian processes on $\mathbb{R}^d$;

Assumption 5. A fixed $\Delta > 0$ exists so that, for all $n \in \mathbb{N}$ and, for all $i \neq j$, $i, j = 1, \dots, n$, $\|s_i - s_j\| \geq \Delta$;

Assumption 6. Fixed $A > 0$ and $\alpha > 0$ exist such that, for all $x \in \mathbb{R}^d$ and, for all $k = 1, \dots, p$,

$$|K_k(x)| \leq \frac{A}{1 + \|x\|^{d+\alpha}};$$

Assumption 7. Assuming Assumption 6 holds, then for the same $A > 0$ and $\alpha > 0$ we have

$$|f(x)| \leq \frac{A}{1 + \|x\|^{d+\alpha}};$$

Assumption 8. We have

$$\liminf_{n \rightarrow \infty} \min_{i=2, \dots, p} \left[\left\{ \Omega^{-1} M(f) \Omega^{-T} \right\}_{i,i} - \left\{ \Omega^{-1} M(f) \Omega^{-T} \right\}_{i-1, i-1} \right] > 0.$$

Assumption 5 implies that $\mathcal{S}^d$ is unbounded as $n \rightarrow \infty$, which means that we address the increasing domain asymptotic framework (Cressie, 1993).

Assumption 7 holds in particular for the function $I(s = 0)$ and for the “ball” and “ring” kernels $B(h)(s) = I(\|s\| \leq h)$ with fixed $h \geq 0$ and $R(h_1, h_2)(s) = I(h_1 \leq \|s\| \leq h_2)$ with fixed $h_2 \geq h_1 \geq 0$.

Up to reordering the components of Z , which comes without loss of generality, Assumption 8 is an asymptotic version of the identifiability condition in Proposition 1. Under Assumption 8, identifiability in the sense of Definition 3 holds for sufficiently large n , from Proposition 1.

Proposition 2 below gives the consistency of the estimator $\widehat{M}(f)$, where f satisfies Assumption 7. The proof of this proposition is provided in § B.4 of the supplementary material.

PROPOSITION 2. *Suppose $n \rightarrow \infty$ and Assumptions 1 to 6 hold and let $f : \mathbb{R}^d \rightarrow \mathbb{R}$ satisfy Assumption 7. Then $\widehat{M}(f) - M(f) \rightarrow 0$ in probability when $n \rightarrow \infty$.*

We remark that $M(f)$ depends on n and that we do not assume that the sequence of matrices $M(f)$ converges to a fixed matrix as $n \rightarrow \infty$. Hence, Proposition 2 shows that $\widehat{M}(f) - M(f)$ converges to zero, and not that $\widehat{M}(f)$ converges to $M(f)$.

Next, we show the joint asymptotic normality of $n^{1/2}\{\widehat{M}(f_0) - M(f_0)\}$ and $n^{1/2}\{\widehat{M}(f) - M(f)\}$, seen as sequences of $p^2 \times 1$ random vectors. Similarly as in Proposition 2, we do not need to assume that the sequence of $2p^2 \times 2p^2$ covariance matrices of these two sequences of vectors converges to a fixed matrix. Hence, we will not show that these sequences of random vectors converge jointly to a fixed Gaussian distribution. Instead, we will show that the distances between the distributions of these random vectors and Gaussian distributions converge to zero as $n \rightarrow \infty$. As a distance between distributions, we consider a metric d_w generating the topology of weak convergence on the set of Borel probability measures on Euclidean spaces (see, e.g., Dudley (2002), p. 393). The benefit of such a distance is that a sequence of distributions $(\mathcal{L}_n)_{n \in \mathbb{N}}$ converges to a fixed distribution $\mathcal{L}$ if and only if $d_w(\mathcal{L}_n, \mathcal{L})$ converges to zero. The next proposition provides the asymptotic normality result. It is proved in § B.4 of the supplementary material.

PROPOSITION 3. *Assume the same assumptions as in Proposition 2. Let $W(f)$ be the vector of size $p^2 \times 1$, defined for $i = (a - 1)p + b$, $a, b \in \{1, \dots, p\}$, by*

$$W(f)_i = n^{1/2}\{\widehat{M}(f)_{a,b} - M(f)_{a,b}\}.$$

Let Q_n be the distribution of $\{W(f)^\top, W(f_0)^\top\}^\top$. Then, as $n \rightarrow \infty$,

$$d_w[Q_n, \mathcal{N}\{0, V(f, f_0)\}] \rightarrow 0,$$

where $\mathcal{N}$ denotes the normal distribution and details concerning the matrix $V(f, f_0)$ are given in Appendix A.2. Furthermore, the largest eigenvalue of $V(f, f_0)$ is bounded as $n \rightarrow \infty$.

In Proposition 3, $V(f, f_0)$ is a $2p^2 \times 2p^2$ matrix that depends on n and is interpreted as an asymptotic covariance matrix. Also, in Proposition 3, the vectors $W(f)$ and $W(f_0)$, that are asymptotically Gaussian, are obtained by row vectorization of $n^{1/2}\{\widehat{M}(f_0) - M(f_0)\}$ and $n^{1/2}\{\widehat{M}(f) - M(f)\}$. Taking $f(s) = I(\|s\| \leq h)$ with $h > 0$ in Propositions 2 and 3 gives the asymptotic properties of the method proposed in Nordhausen et al. (2015).

Remark 1. Propositions 2 and 3 remain valid when centering the process X by $\bar{X} = n^{-1} \sum_{i=1}^n X(s_i)$. Indeed, we prove in Lemma B.8 of the supplementary material that the difference between the centered estimator and $\widehat{M}(f)$ is of order $O_p(n^{-1})$.

For a matrix A with rows $l_1^\top, \dots, l_k^\top$, let $\text{vect}(A) = (l_1^\top, \dots, l_k^\top)^\top$ be the row vectorization of A and for a matrix A of size $k \times k$, let $\text{diag}(A) = (A_{1,1}, \dots, A_{k,k})^\top$. Next, Proposition 4 shows the joint asymptotic normality of the estimators $\widehat{\Gamma}(f)$ and $\widehat{\Lambda}(f)$. This proposition is proved in § B.4 of the supplementary material.

PROPOSITION 4. Assume the same assumptions as in Proposition 2. Assume also that Assumption 8 holds. For $\widehat{\Gamma}(f)$ and $\widehat{\Lambda}(f)$ in Definition 1, let Q_n be the distribution of

$$n^{1/2} \begin{pmatrix} \text{vect} \left\{ \widehat{\Gamma}(f) - \Omega^{-1} \right\} \\ \text{diag} \left\{ \widehat{\Lambda}(f) - \Lambda(f) \right\} \end{pmatrix}.$$

Then, we can choose $\widehat{\Gamma}(f)$ and $\widehat{\Lambda}(f)$ in Definition 1 so that when $n \rightarrow \infty$,

$$d_w\{Q_n, \mathcal{N}(0, F_1)\} \rightarrow 0,$$

where details concerning the matrix F_1 are given in Appendix A.3.

In Proposition 4, similarly as before, we consider the sequences of vectors obtained by vectorizing $n^{1/2}\{\widehat{\Gamma}(f) - \Omega^{-1}\}$ and taking the diagonal of $n^{1/2}\{\widehat{\Lambda}(f) - \Lambda(f)\}$. Again, we do not show that the sequence of joint distributions of these vectors converges to a fixed distribution. Instead, we show that these joint distributions are asymptotically close to Gaussian distributions, with covariance matrices given by F_1 . We remark that F_1 denotes a sequence of $(p^2 + p) \times (p^2 + p)$ matrices. We also remark that, in Definition 1, $\widehat{\Gamma}(f)$ is not uniquely defined. It is defined up to the signs of its rows. Hence, Proposition 4 shows that there exists a choice of the sequence $\widehat{\Gamma}(f)$ in Definition 1 such that asymptotic normality holds as $n \rightarrow \infty$.

The performance of the estimators $\widehat{\Gamma}(f)$ and $\widehat{\Lambda}(f)$ depends on the choice of $\widehat{M}(f)$ that should be chosen so that $\widehat{\Lambda}(f)$ has diagonal elements as distinct as possible. This is similar to the time series context as described in Miettinen et al. (2012). To avoid this dependency in the time series context, the joint diagonalisation of more than two matrices has been suggested and we will apply this concept to the spatial context in the following section.

4. IMPROVING THE ESTIMATION OF THE SPATIAL BLIND SOURCE SEPARATION MODEL BY JOINTLY DIAGONALISING MORE THAN TWO MATRICES

Spatial blind source separation with more than two kernel functions of the form $f_0, f_1, \dots, f_k$, with $k \geq 2$, can be formulated as

$$\widehat{\Gamma} \in \underset{\substack{\Gamma: \Gamma \widehat{M}(f_0) \Gamma^T = I_p \\ \Gamma \text{ has rows } \gamma_1^T, \dots, \gamma_p^T}}{\text{argmax}} \sum_{l=1}^k \sum_{j=1}^p \{\gamma_j^T \widehat{M}(f_l) \gamma_j\}^2. \quad (5)$$

We can show that, if $k = 1$, the set of $\widehat{\Gamma}$ satisfying (5) coincides with the set of $\widehat{\Gamma}(f_1)$ satisfying Definition 1. From experience in time series blind source separation (see for example Miettinen et al., 2016), usually the diagonalisation of several matrices gives a better separation than those based on two matrices only. In this paper, we indeed show that using $k \geq 2$ is beneficial from a theoretical point of view and in practice.

The identifiability notion of Definition 3 and Proposition 1 can be extended to the case of more than two local covariance matrices. We first remark that the theoretical version of (5) is

$$\Gamma \in \underset{\substack{\Gamma: \Gamma M(f_0) \Gamma^T = I_p \\ \Gamma \text{ has rows } \gamma_1^T, \dots, \gamma_p^T}}{\text{argmax}} \sum_{l=1}^k \sum_{j=1}^p \{\gamma_j^T M(f_l) \gamma_j\}^2. \quad (6)$$

We then extend Definition 3 and Proposition 1 to the case of more than two local covariance matrices.

DEFINITION 4. *We say that the unmixing problem given by $f_1, \dots, f_k$ is identifiable if any unmixing functional Γ satisfying (6) can be written as $PS\Omega^{-1}$, where P is a permutation matrix and S is a diagonal matrix with diagonal elements equal to -1 or 1 .*

PROPOSITION 5. *The unmixing problem given by $f_1, \dots, f_k$ is identifiable if and only if for every pair $i \neq j$, $i, j = 1, \dots, p$, there exists $l = 1, \dots, k$ such that $\{\Omega^{-1}M(f_l)\Omega^{-T}\}_{i,i} \neq \{\Omega^{-1}M(f_l)\Omega^{-T}\}_{j,j}$.*

Proposition 5 is proved in § B.5 of the supplementary material. We remark that the identifiability condition in Proposition 5 is weaker than that in Proposition 1, because if the condition in Proposition 1 holds with f being one of the $f_1, \dots, f_k$, then the condition in Proposition 5 holds. This is one of the benefits of jointly diagonalising more than two matrices.

One of the main theoretical contributions of this paper is to provide an asymptotic analysis of the joint diagonalisation of several matrices in the spatial context. Assumption 8, on asymptotic identifiability, can be replaced by the following weaker assumption.

Assumption 9. A fixed $\delta > 0$ and $n_0 \in \mathbb{N}$ exist so that for all $n \in \mathbb{N}$, $n \geq n_0$, for every pair $i \neq j$, $i, j = 1, \dots, p$, there exists $l = 1, \dots, k$, such that $|\{\Omega^{-1}M(f_l)\Omega^{-T}\}_{i,i} - \{\Omega^{-1}M(f_l)\Omega^{-T}\}_{j,j}| \geq \delta$.

In the next proposition, we prove the consistency of $\hat{\Gamma}$. This proposition is proved in § B.6 of the supplementary material.

PROPOSITION 6. *Suppose Assumptions 1 to 6 hold. Let $k \in \mathbb{N}$ be fixed. Let $f_1, \dots, f_k : \mathbb{R}^d \rightarrow \mathbb{R}$ satisfy Assumption 7. Assume that Assumption 9 holds. Let $\hat{\Gamma} = \hat{\Gamma}\{\hat{M}(f_0), \hat{M}(f_1), \dots, \hat{M}(f_k)\}$ satisfy (5). Then we can choose $\hat{\Gamma}$ so that $\hat{\Gamma} \rightarrow \Omega^{-1}$ in probability when n goes to infinity.*

In Proposition 6, we remark that $\hat{\Gamma}$ is defined only up to permutation of the rows and multiplications of them by 1 or -1 . Hence, we show that there exists a choice of a sequence $\hat{\Gamma}$ that converges to Ω^{-1} . The next proposition provides an asymptotic normality result. It is proved in § B.6 of the supplementary material.

PROPOSITION 7. *Assume the same assumptions as in Proposition 6. Let $(\hat{\Gamma}_n)_{n \in \mathbb{N}}$ be any sequence of $p \times p$ matrices so that for any $n \in \mathbb{N}$, $\hat{\Gamma}_n = \hat{\Gamma}_n\{\hat{M}(f_0), \hat{M}(f_1), \dots, \hat{M}(f_k)\}$ satisfies (5). Then, a sequence of permutation matrices (P_n) and a sequence of diagonal matrices (D_n) exist, with diagonal components in $\{-1, 1\}$, so that the distribution Q_n of $n^{1/2}\text{vect}(\check{\Gamma}_n - \Omega^{-1})$ with $\check{\Gamma}_n = D_n P_n \hat{\Gamma}_n$ satisfies, as $n \rightarrow \infty$,*

$$d_w\{Q_n, \mathcal{N}(0, F_k)\} \rightarrow 0,$$

where details concerning the matrix F_k are given in Appendix A.4.

In Proposition 7, for any $n \in \mathbb{N}$, the choice of $\hat{\Gamma}_n$ satisfying (5) is not unique. The proposition shows that, for any choice of the sequence of matrices $\hat{\Gamma}_n$, one can exchange the rows and multiply them by 1 or -1 , to obtain a sequence of matrices $\check{\Gamma}_n$ that converges to Ω^{-1} as $n \rightarrow \infty$. Furthermore, similarly as in Proposition 4, we show that the sequence of distributions of $n^{1/2}\text{vect}(\check{\Gamma}_n - \Omega^{-1})$ is asymptotically close to a sequence of Gaussian distributions. The sequence of $p^2 \times p^2$ covariance matrices of these Gaussian distributions is F_k .

The idea of joint diagonalisation is not new in spatial data analysis. For example in Xie & Myers (1995), Xie et al. (1995) and De Iaco et al. (2013), in a model-free context, matrix variograms have been jointly diagonalized. However, the unmixing matrix was restricted to be orthogonal, which would therefore not solve the spatial blind source separation model.

While two symmetric matrices can always be simultaneously diagonalized, this is usually not the case for more than two matrices which are estimated based on finite data. Therefore, algorithms are needed for approximate joint diagonalisation. In this paper we use an algorithm which is based on Givens rotations (Clarkson, 1988). Other possible algorithms and their impact on the properties of the estimates are for example discussed in Illner et al. (2015).

5. SIMULATIONS

5.1. Preliminaries

In this section we use simulated data to verify our asymptotic results and to compare the efficiencies of the different local covariance estimates under a varying set of spatial models. All simulations are performed in R (R Core Team, 2019) with the help of the packages *geoR* (Ribeiro Jr & Diggle, 2016), *JADE* (Miettinen et al., 2017) and *RcppArmadillo* (Eddelbuettel & Sanderson, 2014). To generate the simulation data, we have chosen some particular covariance functions for the latent fields. However, our proposed methods do not use this information in any way, but operate solely through the selection of local covariance matrices.

5.2. Asymptotic approximation of the unmixing matrix estimator

We start with a simple simulation to establish the validity of the asymptotic approximation of the unmixing matrix estimator $\hat{\Gamma}(f)$ for different kernels f and to obtain some preliminary comparative results between the proposed estimators. We consider a centered, three-variate spatial blind source separation model $X(s) = \Omega Z(s)$ where each of the three independent latent fields has a Matérn covariance function with shape and range parameters $(\kappa, \phi) \in \{(6, 1.2), (1, 1.5), (0.25, 1)\}$, which correspond to the left panel in Fig. 1. We recall that the Matérn correlation function is defined by

$$\rho(h) = 2^{1-\kappa} \Gamma(\kappa)^{-1} (h/\phi)^\kappa K_\kappa(h/\phi),$$

where $\kappa > 0$ is the shape parameter, $\phi > 0$ is the range parameter and K_κ is the modified Bessel function of the second kind of order κ . Our location pattern is constructed in the following way: the first 200 locations are drawn uniformly random from an origin-centered square S_1 of side length $200^{1/2}$ units. For the next 200 locations, we scale the side length of the square S_1 by the factor $2^{1/2}$ to obtain the larger square S_2 and draw the points uniformly random on $S_2 \setminus S_1$. Next, we always scale the side length of the previous square S_j by $2^{1/2}$ to obtain S_{j+1} and draw the same amount of locations we already have on $S_{j+1} \setminus S_j$, thus doubling the number of points every time. This process is continued until we have obtained a total of 3200 locations. In the simulation we consider the sample sizes $n = 100 \times 2^j$, for $j = 1, \dots, 5$, each time using the first n of the 3200 points, that is, all points inside the j th innermost square on the left-hand side of Fig. 2. The six samples then correspond to nested samples of points and represent the increasing domain asymptotic scheme implied by Assumption 5.

We expect any successful unmixing estimator $\hat{\Gamma}$ to satisfy $\hat{\Gamma}\Omega \approx I_p$ up to sign changes and row permutations. The minimum distance index (Ilmonen et al., 2010b) is defined as,

$$\text{MDI}(\hat{\Gamma}) = (p-1)^{-1/2} \inf\{\|C\hat{\Gamma}\Omega - I_p\|, C \in \mathcal{C}\},$$

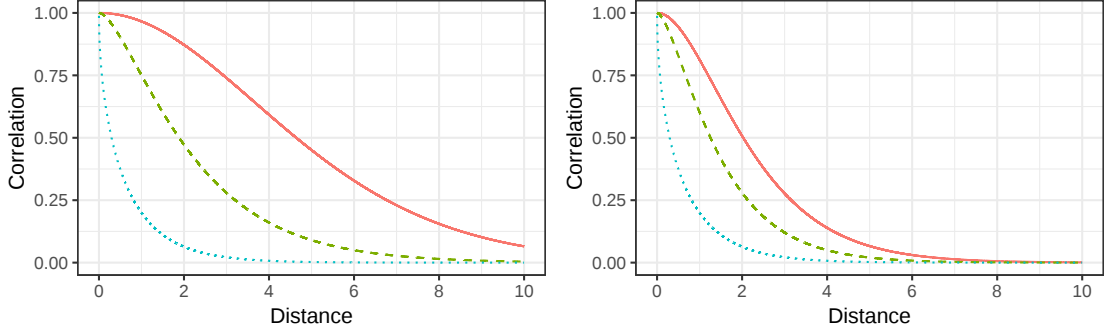


Fig. 1: Matérn covariance functions of the first (solid red line), second (dashed green line) and third (dotted blue line) latent fields used in § 5.2 (left panel) and § 5.3 (right panel). The parameter vectors (κ, ϕ) of the three fields equal, $(6, 1.2)$, $(1, 1.5)$, $(0.25, 1)$ and $(2, 1)$, $(1, 1)$, $(0.25, 1)$, respectively.

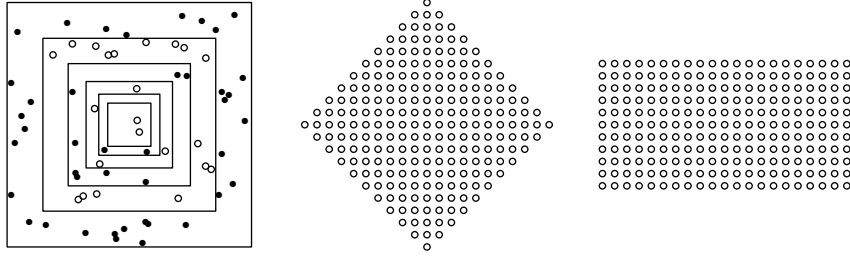


Fig. 2: From left to right: the location pattern scheme used in § 5.2 (with the marker type alternating between the consecutive layers and only 1 percent of the locations shown for clarity), a diamond grid of radius 10 having $n = 221$ locations and a rectangle grid of radius 10 having $n = 231$ locations. The diamond and rectangle grids, with a one unit distance between two neighbouring locations, are used in § 5.3.

where $\mathcal{C}$ is the set of all matrices with exactly one non-zero element in each row and column and $\|\cdot\|$ is the Frobenius norm. The minimum distance index measures how close $\hat{\Gamma}\Omega$ is to the identity matrix up to scaling, order and signs of its rows, and $0 \leq \text{MDI}(\hat{\Gamma}) \leq 1$ with lower values indicating more efficient estimation. Moreover, for any $\hat{\Gamma}$ such that $n^{1/2} \text{vect}(\hat{\Gamma} - I_p) \rightarrow \mathcal{N}(0, \Sigma)$ for some limiting covariance matrix Σ , the transformed index $n(p-1)\text{MDI}(\hat{\Gamma})^2$ converges to a limiting distribution $\sum_{i=1}^k \delta_i \chi_i^2$ where $\chi_1^2, \dots, \chi_k^2$ are independent chi-squared random variables with one degree of freedom and $\delta_1, \dots, \delta_k$ are the k non-zero eigenvalues of the matrix,

$$(I_{p^2} - D_{p,p}) \Sigma (I_{p^2} - D_{p,p}),$$

where $D_{p,p} = \sum_{j=1}^p E^{jj} \otimes E^{jj}$ and E^{jj} is the $p \times p$ matrix with one as its (j, j) th element and the rest of the elements equal zero, and $\otimes$ is the usual tensor matrix product. In particular, the expected value of the limiting distribution is the sum of the limiting variances of the off-diagonal elements of $\hat{\Gamma}$. This provides us with a useful single-number summary to measure the asymptotic efficiency of the method, i.e., the mean value of $n(p-1)\text{MDI}(\hat{\Gamma})^2$ over several replications.

Following the argument of the proof of Proposition B.4 in the supplementary material, our spatial blind source separation estimators are affine equivariant. More precisely, let $\hat{\Gamma}(I_p)$ be computed from $\{Z(s_i)\}_{i=1, \dots, n}$ according to (5) and recall that $\hat{\Gamma}$ is computed from $\{X(s_i)\}_{i=1, \dots, n}$

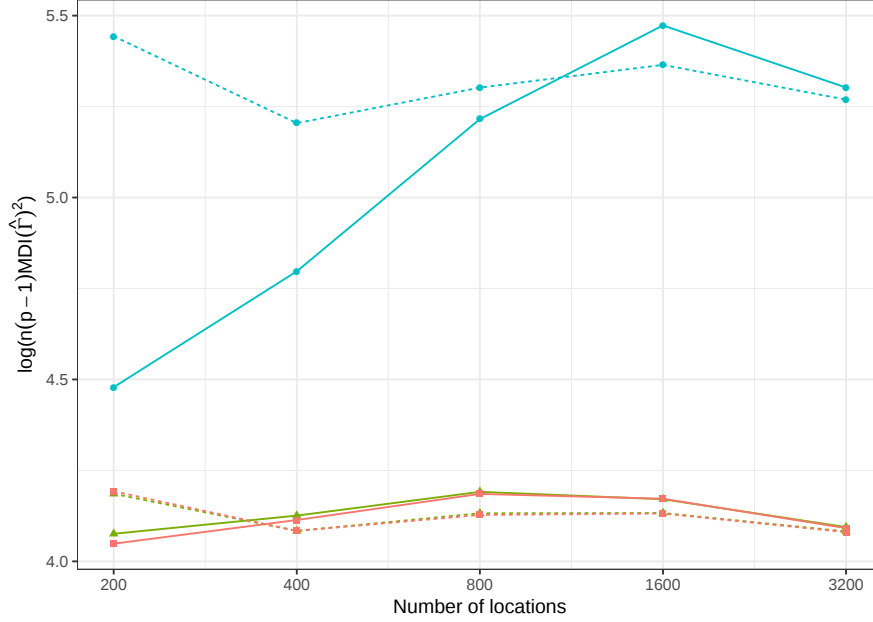


Fig. 3: The solid lines give the mean values of $n(p-1)\text{MDI}(\hat{\Gamma})^2$ in the first simulation and the dashed lines correspond to the asymptotic approximations of the same quantities. The three used local covariance matrices are $B(1)$ (blue line), $R(1, 2)$ (green line) and $\{B(1), R(1, 2)\}$ (red line).

according to (5). Then we have $\hat{\Gamma} = \hat{\Gamma}(I_p)\Omega^{-1}$, up to sign changes and row permutations. In this sense, $\hat{\Gamma}\Omega$ is invariant to the value of Ω . As the minimum distance index depends on $\hat{\Gamma}$ only through $\hat{\Gamma}\Omega$, it is thus without loss of generality that we may consider throughout § 5 only the trivial mixing case $\Omega = I_3$. Taking different Ω into consideration would give exactly the same results as those provided below.

Recall that the ball and ring kernels are defined as $B(h)(s) = I(\|s\| \leq h)$ and $R(h_1, h_2)(s) = I(h_1 \leq \|s\| \leq h_2)$ for fixed $h \geq 0$ and $h_2 \geq h_1 \geq 0$. We simulate 2000 replications for each sample size n and estimate the unmixing matrix in each case with three different choices for the local covariance matrix kernels: $B(1)$, $R(1, 2)$ and $\{B(1), R(1, 2)\}$, where the argument s is dropped and the brackets $\{\}$ denote the joint diagonalisation of the kernels inside. The latent covariance functions on the left panel of Fig. 1 show that the dependencies of the last two fields die off rather quickly, and we would expect that already very local information is sufficient to separate the fields. Moreover, out of all one-unit intervals, the magnitudes of the three covariance functions differ the most from each other in the interval from 1 to 2 and we may reasonably assume that either $R(1, 2)$ or $\{B(1), R(1, 2)\}$ will be the most efficient choice.

The mean values of $n(p-1)\text{MDI}(\hat{\Gamma})^2$ over the 2000 replications are shown as the solid lines in Fig. 3, with the dashed lines representing the asymptotic approximated values of the means, towards which they are expected to converge (see Propositions 4 and 7). As evidenced in Fig. 3, this is indeed what happens. For the reasons detailed in the previous paragraph, the kernel $R(1, 2)$ is notably a more efficient choice than $B(1)$. However, the ball kernel still carries some additional information to the ring as their joint diagonalisation, $\{B(1), R(1, 2)\}$, gives the best results out of the three choices, albeit marginally. As the main purpose of the current simulation was to verify the limiting theorems and compare the different choices of kernels, the estimation accuracy of the

sources was considered jointly, through the minimum distance index. However, as it is possible that some of the individual sources are more difficult to estimate than others, we have included in § C.2 of the supplementary material a simulation study exploring individual component recovery.

The previous investigation and Fig. 3 used only the expected value of the asymptotic distribution. In Fig. C.1 of the supplementary material, we have also plotted the estimated densities of $n(p-1)\text{MDI}(\hat{\Gamma})^2$ for all local covariance matrices and a few selected sample sizes and compared with the density of the asymptotic approximation estimated from a sample of 100,000 random variables drawn from the corresponding distributions. Overall, the two densities fit each other rather well, especially for the local covariance matrices involving the ring kernel. This shows that the asymptotic approximation to the distribution of $n(p-1)\text{MDI}(\hat{\Gamma})^2$ is good already for small sample sizes.

5.3. The effect of range on the efficiency

The second simulation explores the effect of the range of the latent fields on the asymptotically optimal choice of local covariance matrices. The comparisons between the estimators are made on the basis of the expected values of the asymptotic approximations to the distribution of $n(p-1)\text{MDI}(\hat{\Gamma})^2$ (that is, using the equivalent of the dashed lines in Fig. 3), meaning that no randomness is involved in this simulation.

We consider three-variate random fields $X(s) = \Omega Z(s)$, where $\Omega = I_3$ and the latent fields have Matérn covariance functions with respective shape parameters $\kappa = 2, 1, 0.25$ and a range parameter $\phi \in \{1.0, 1.1, 1.2, \dots, 30.0\}$. The three covariance functions are shown for $\phi = 1$ in the right panel of Fig. 1. The random field is observed at three different point patterns: diamond-shaped, rectangular and random, which was simulated once and held fixed throughout the study. The diamond-shaped point pattern has a radius of $m = 30$ and a total of $n = 1861$ locations, whereas the rectangular point pattern has a “radius” of $m = 15$ with a total of $n = 1891$ locations. In both patterns, the horizontal and vertical distance between two neighbouring locations is one unit and examples of the two pattern types are shown in the middle and right panels of Fig. 2 with a radius $m = 10$. A rectangular pattern with “radius” m is defined to have the width $2m + 1$ and the height $m + 1$. The random point pattern is generated simply by simulating $n = 1861$ points uniformly in the rectangle $(-30, 30) \times (-30, 30)$. We consider a total of eight different local covariance matrices, $B(r)$, $R(r-1, r)$ for $r = 1, 3, 5$, and the joint diagonalisations of the previous sets: $\{B(1), B(3), B(5)\}$ and $\{R(0, 1), R(2, 3), R(4, 5)\}$.

The results of the simulation are displayed in Fig. 4 where the two joint diagonalisations are denoted by having value “J” as the parameter r . Recall that the lower the value on the y -axis, the better that particular method is at estimating the three latent fields. The relative ordering of the different curves is very similar across all three plots, and it seems that the choice of the location pattern does not have a large effect on the results. In all the patterns, the local covariance matrices with either $r = 1$ or $r = 3$ are the best choices for small values of the range ϕ but they quickly deteriorate as ϕ increases. The opposite happens for the local covariance matrices with $r = 5$; they are among the worst for small ϕ and relatively improve with increasing ϕ . The joint diagonalisation-based choices fall somewhere in-between and are never the best nor the worst choice. However, they yield performance very close to the best choice in the right end of the range-scale and are close to the optimal ones in the left end. Thus, their use could be justified in practice as the “safe choice”. Comparing the two types of local covariance matrices, balls and rings, we observe that in the majority of cases the rings prove superior to the balls.

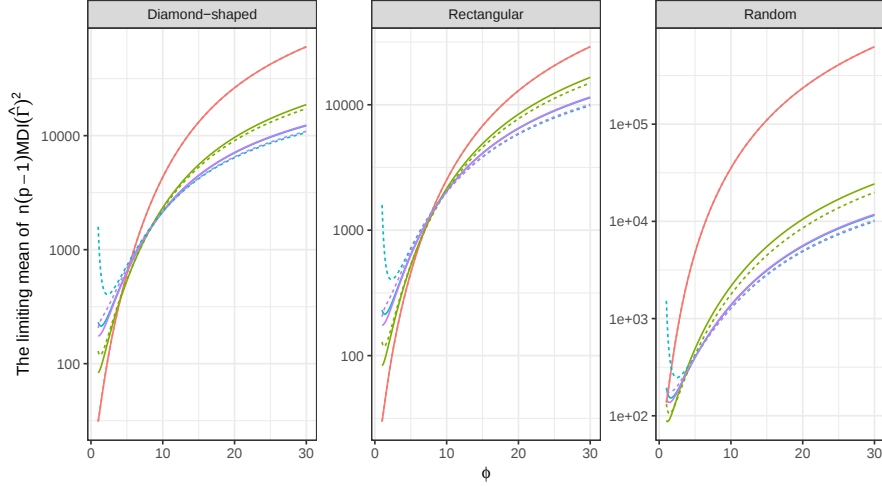


Fig. 4: The asymptotic approximate mean values of $n(p-1)\text{MDI}(\hat{\Gamma})^2$ as a function of the range of the latent Matérn random fields for the different choices of local covariance matrices in the second simulation. The solid and the dashed lines correspond, respectively, to the ball and ring kernels and the value of the parameter r is indicated by the color of the line as follows: 1 (red), 3 (green), 5 (blue), J (purple). The y -axis has a logarithmic scale.

5.4. Efficiency comparison

To compare a larger number of local covariance matrices and their combinations, we simulate three-variate random fields $X(s) = \Omega Z(s)$, where $\Omega = I_3$ and the latent fields have Matérn covariance functions with the shape parameters $\kappa = 6, 1, 0.25$ and the range parameter $\phi = 20$, in kilometers. We consider two different fixed-location patterns fitted inside the map of Finland; see Fig. 5. The first location pattern has the locations drawn uniformly from the map and the second location pattern is drawn from a west-skew distribution. Both patterns have a total of $n = 1000$ locations and to better distinguish the scale we have added three concentric circles with respective radii of 10, 20, and 30 kilometers in the empty area of the skew map.

We simulate a total of 2000 replications of the above scheme with the fixed maps. In each case we compute the minimum distance index values of the estimates obtained with the local covariance matrix kernels $B(r), R(r-10, r), G(r)$, where $r = 10, 20, 30, 100$, and the joint diagonalisation of each of the three quadruplets $\{B(10), B(20), B(30), B(100)\}$, $\{R(10), R(20), R(30), R(100)\}$ and $\{G(10), G(20), G(30), G(100)\}$ adding up to a total of 15 estimators. The Gaussian kernel is parametrized as $G(r) \equiv \exp[-0.5\{\Phi^{-1}(0.95)s/r\}^2]$, where s is the distance and $\Phi^{-1}(x)$ is the quantile function of the standard normal distribution, making $G(r)$ have 90% of its total mass in the radius r ball around its center. Thus, $G(r)$ can be considered a smooth approximation of $B(r)$. The larger radius kernels $B(100), R(90, 100), G(100)$ are included in the simulation to investigate what happens when we overestimate the dependency radius. The mean minimum distance index values for the 15 estimators are plotted in Fig. 6 and show that for both maps and all local covariance types, increasing the radius yields more accurate separation results all the way up to $r = 30$, whereas for $r = 100$ the results again worsen. This observation shows that when using a single local covariance matrix, the choice of the type and the radius are especially important, most likely requiring some expert knowledge on the study. However, this problem is completely averted when we use the joint diagonalisation of several matrices. For both maps and all local covariance types the joint diagonalisation produces results very comparable to the best individual matrices, even though the joint diagonalisations

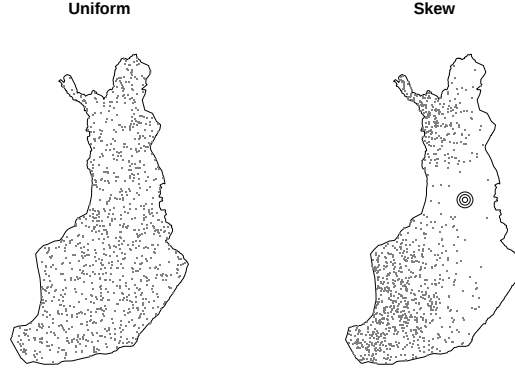


Fig. 5: The two fixed location patterns in the map of Finland, the uniform on the left and the skew on the right.

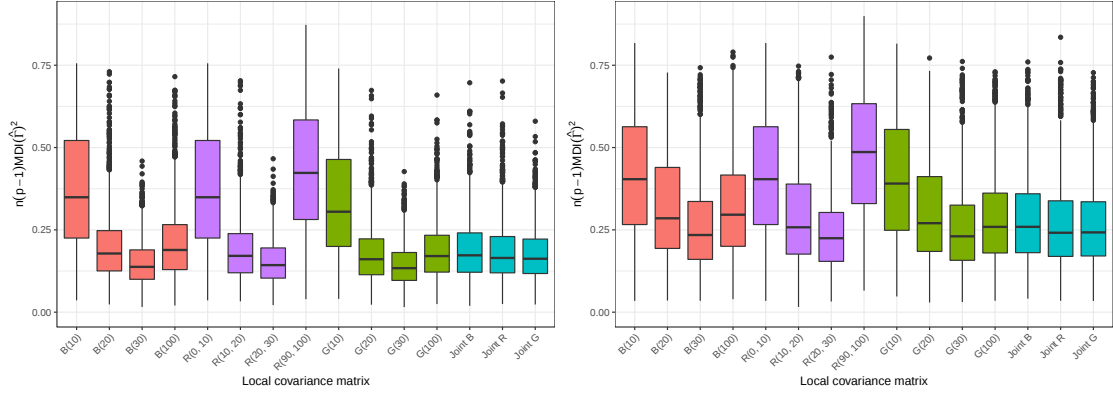


Fig. 6: The results of the efficiency study for the uniform sampling design (left) and for the skew design (right).

also include the “bad choices”, $r = 10, 20, 100$. We also observe a similar behaviour in the first and second simulation studies where, in the absence of knowledge on the optimal choice, the joint diagonalisation either is the most efficient choice or provides a performance very close to the most efficient choice. Thus, we recommend the use of the joint diagonalisation of scatter matrices with a sufficiently large variation of radii for the kernels.

Finally, a comparison between the two maps reveals that the relative behaviour of the estimators is roughly the same in both maps, but the estimation is generally more difficult in the skew map, revealed by the on average higher minimum distance index values. This is explained by the large number of isolated points which contribute no information to the estimation of the local covariance matrices, making the sample size essentially smaller than $n = 1000$.

6. DATA APPLICATION

To illustrate the benefit of jointly diagonalising more than two scatter matrices from a practical point of view, we reconsider the moss data from the Kola project which are available in the R package StatDa (Filzmoser, 2015) and described in Reimann et al. (2008), for example. The data consist of 594 samples of terrestrial moss collected at different sites in north Europe on the borders of Norway, Finland and Russia. The corresponding map with sampling locations

Table 1: Maximal absolute correlations of different estimators with respect to the gold standard. All estimators used the empirical covariance matrix. The distances for the scatters are given in kilometers

Est	Scatters	IC1	IC2	IC3	IC4	IC5	IC6
1	$B(25)$	0.96	0.93	0.91	0.68	0.64	0.77
2	$B(75)$	0.98	0.98	0.92	0.96	0.91	0.63
3	$B(100)$	0.76	0.80	0.77	0.96	0.60	0.53
4	$R(0, 25), R(25, 50), R(50, 75), R(75, 100)$	0.97	0.98	0.92	0.97	0.83	0.80
5	$R(0, 10), R(10, 20), R(20, 30), R(30, 40),$ $R(40, 50), R(50, 60), R(60, 70), R(70, 80)$	0.96	0.97	0.91	0.97	0.78	0.77

Est, estimator; IC, independent component.

is given in the online supplement in Fig. D.1. The amount of 31 chemical elements found in the moss samples was already used as a spatial blind source separation example in Nordhausen et al. (2015) where the covariance matrix and $B(50)$ were simultaneously diagonalized. The goal of that analysis was to reveal interpretable components exhibiting clear spatial patterns. In Nordhausen et al. (2015), the radius of 50 kilometers was carefully chosen by an expert in that analysis and considered best compared to several other radii not mentioned there. The analysis found six meaningful components, which could be used to distinguish underlying natural geological patterns from environmental pollution patterns. These six components had the six largest eigenvalues and are visualized in Fig. D.2 in the online supplement.

We show that the gold standard components can be stably estimated without subject knowledge on the optimal radius by simply jointly diagonalizing a large enough collection of local covariance matrices. To address the compositional nature of the data, we follow the same data preparation steps as in Nordhausen et al. (2015) and then compute five competing spatial blind source separation estimates. The scatters we used in addition to the covariance matrix are detailed in Table 1. Using these methods, we identify the six components with the highest correlations, in absolute values, to the six main components identified in Nordhausen et al. (2015). Table 1 gives the correlations of the six components. The table shows that when using only two scatters, estimators 1, 2 and 3, some components cannot be easily found. However, when jointly diagonalising more than two scatters, the results are more stable and less dependent on the chosen distances of the scatters as can be seen for estimators 4 and 5.

This is illustrated using the gold standard and estimators 3 and 4 in Fig. A.1 in the Appendix for the first two components. For completeness, § D of the online supplement contains all six components for the three estimators. The first two components represent, according to Nordhausen et al. (2015), areas with different types of industrial contamination and Figure A.1 shows that the gold standard and estimator 4 agree quite well on these, but estimator 3 yields a different map. More precisely, the first component obtained by the gold standard and the estimator 4 highlights a cluster of negative scores around the Monchegorsk and Apatity region, which reveals the mining and processing of alkaline deposits. This cluster is not revealed by estimator 3. Similarly, the second components are similar between the gold standard and the estimator 4, but the component from the estimator 3 differs from these two, especially for the sampling locations in Finland. Thus, using several scatters gives a more stable impression whereas the maps can vary considerably when only two scatters are used, in which case subject expertise becomes more relevant.

7. DISCUSSION

Our proposed methodology can be extended in multiple directions in future work. The assumptions of Gaussian or stationary fields could be relaxed. The spatial and temporal blind source separation methodologies could be combined to obtain spatio-temporal blind source separation. If used for dimension reduction, estimators for the number of latent non-noise fields could be devised using strategies similar to those in Virta & Nordhausen (2019). Additionally, the combination of spatial blind source separation with univariate kriging and univariate modelling warrants investigation.

How to choose the local covariance matrices optimally is also of interest. This is still an open problem for temporal blind source separation methods, such as second-order blind identification (Belouchrani et al., 1997). Several strategies have been suggested, see for example Tang et al. (2005), and many of them could be useful also in selecting the kernels in spatial blind source separation. The estimation accuracy of our proposed method is based on how well separated the eigenvalues of the matrices $M(f_0)^{-1/2}M(f_l)M(f_0)^{-1/2}$, $l = 1, \dots, k$, are. Since the connection between the eigenvalues and the unknown covariance functions is complicated, our suggestion, backed up also by the simulations, is to stay on the safe side and jointly use a large number of ring kernels. However, including large numbers of unnecessary kernels can still have the drawback of inducing some noise to the estimates. One way to remove the unneeded kernels would be to first obtain preliminary estimates for the latent fields using a large number of kernels jointly. Then, our asymptotic results could be used to select from a large collection of sets of kernels, the one which achieves the smallest value of $\delta_1 + \dots + \delta_k$; see § 5.2. The final estimates could then be computed with this asymptotically optimal choice of kernels. A similar technique was used in the context of temporal blind source separation in Taskinen et al. (2016).

ACKNOWLEDGEMENT

The work of Nordhausen, Ruiz-Gazen and Virta was partly supported by the CRoNoS Cost action. The work of Nordhausen was also partly supported by the Austrian Science Fund. Ruiz-Gazen acknowledges funding from the French National Research Agency (ANR) under the Investments for the Future (Investissements d'Avenir) program. The authors are very grateful for the comments by the referees which helped considerably to improve the manuscript.

A. APPENDIX

A.1. Notation

Let y and z be the $np \times 1$ vectors defined by $y_{(i-1)p+j} = Y_j(s_i)$ and $z_{(i-1)p+j} = Z_j(s_i)$, for $i = 1, \dots, n$, $j = 1, \dots, p$. Let $R = \text{cov}(y)$ and $R_z = \text{cov}(z)$. Let $e_b(p)$ be the b th base column vector of $\mathbb{R}^p$ for $b = 1, \dots, p$. For $f : \mathbb{R}^d \rightarrow \mathbb{R}$ and for $b, l = 1, \dots, p$, let $T_{b,l}(f)$ be the $np \times np$ matrix, that we see as a block matrix composed of n^2 blocks of sizes p^2 , and with block i, j equal to $f(s_i - s_j)(1/2)\{e_b(p)e_l(p)^T + e_l(p)e_b(p)^T\}$.

For $b \in \mathbb{N}$, we let $\mathcal{D}(b) = \{1 + (i-1)(b+1); i = 1, \dots, b\}$. We remark that $\{\text{vect}(M)_i; i \in \mathcal{D}(b)\} = \{M_{i,i}; i = 1, \dots, b\}$ for a $b \times b$ matrix M . Let $\bar{\mathcal{D}}_b = \{1, \dots, b^2\} \setminus \mathcal{D}_b$. We remark that $\{\text{vect}(M)_i; i \in \bar{\mathcal{D}}(b)\} = \{M_{i,j}; i, j = 1, \dots, b, i \neq j\}$ for a $b \times b$ matrix M . For $a \in \{1, \dots, b^2\}$, let $I_b(a)$ and $J_b(a)$ be the unique $i, j \in \{1, \dots, b\}$ so that $a = b(i-1) + j$. For $i \in \{1, \dots, b\}$, let $d_b(i) = 1 + (i-1)(b+1)$ and note that $\{\text{vect}(M)_{d_b(i)}; i = 1, \dots, b\} = \{M_{i,i}; i = 1, \dots, b\}$ for a $b \times b$ matrix M . For a matrix M of size $b \times b$, recall that $\text{diag}(M) = (M_{1,1}, \dots, M_{b,b})^T$ and that $\text{tr}(M)$ denotes its trace.

A.2. Expression of the matrix $V(f, f_0)$ from Proposition 3

Let $f, g : \mathbb{R}^d \rightarrow \mathbb{R}$. Using the notation of Appendix A.1, let $\Sigma(f)$ and $\Sigma(f, g)$ be the $p^2 \times p^2$ matrices defined by, for $i = (s-1)p + t$ and $j = (u-1)p + v$, with $s, t, u, v \in \{1, \dots, p\}$,

$$\Sigma(f)_{i,j} = 2n^{-1} \text{tr} \{RT(f)_{s,t} RT(f)_{u,v}\} \text{ and } \Sigma(f, g)_{i,j} = 2n^{-1} \text{tr} \{RT(f)_{s,t} RT(g)_{u,v}\}.$$

Let

$$V(f, g) = \begin{pmatrix} \Sigma(f) & \Sigma(f, g) \\ \Sigma(g, f) & \Sigma(g) \end{pmatrix}.$$

Then $V(f, f_0)$ is equal to $V(f, g)$ for $g = f_0$.

A.3. Expression of the matrix F_1 from Proposition 4

From Assumption 8, there exists $n_0 \in \mathbb{N}$ such that for $n \geq n_0$ the diagonal elements of $\Omega^{-1}M(f)\Omega^{-\top}$ are strictly decreasing. Write these diagonal elements as $\lambda_1 > \dots > \lambda_p$. Using the notation of Appendix A.1, for $n \geq n_0$, let A, B, C and D be respectively the $p^2 \times p^2$, $p^2 \times p^2$, $p \times p^2$ and $p \times p^2$ matrices defined by

$$A_{i,j} = \begin{cases} -1/2 & \text{for } i = j \in \mathcal{D}(p), \\ -\lambda_{I_p(i)} \{\lambda_{I_p(i)} - \lambda_{J_p(i)}\}^{-1} & \text{for } i = j \notin \mathcal{D}(p), \\ 0 & \text{otherwise,} \end{cases}$$

$$B_{i,j} = \begin{cases} \{\lambda_{I_p(i)} - \lambda_{J_p(i)}\}^{-1} & \text{for } i = j \notin \mathcal{D}(p), \\ 0 & \text{otherwise,} \end{cases}$$

$$C_{i,j} = \begin{cases} -\lambda_i & \text{for } j = d_p(i), \\ 0 & \text{otherwise} \end{cases} \quad \text{and} \quad D_{i,j} = \begin{cases} 1 & \text{for } j = d_p(i), \\ 0 & \text{otherwise.} \end{cases}$$

Let

$$G = \begin{pmatrix} A & B \\ C & D \end{pmatrix}.$$

Let $M_{\Omega^{-1}}$ and $\bar{M}_{\Omega^{-1}}$ be respectively the $p^2 \times p^2$ and $(p^2 + p) \times (p^2 + p)$ matrices defined by

$$(M_{\Omega^{-1}})_{a,b} = \begin{cases} (\Omega^{-1})_{J_p(b), J_p(a)} & \text{if } I_p(a) = I_p(b), \\ 0 & \text{if } I_p(a) \neq I_p(b) \end{cases} \quad \text{and} \quad \bar{M}_{\Omega^{-1}} = \begin{pmatrix} M_{\Omega^{-1}} & 0 \\ 0 & I_p \end{pmatrix}.$$

Let $\tilde{V}(f)$ be defined as $V(f_0, f)$ but with R replaced by R_z . Then, for $n \geq n_0$, F_1 is defined as

$$F_1 = \bar{M}_{\Omega^{-1}} G \tilde{V}(f) G^\top \bar{M}_{\Omega^{-1}}^\top.$$

A.4. Expression of the matrix F_k from Proposition 7

Let $D(f) = \Omega^{-1}M(f)\Omega^{-\top}$. For a diagonal matrix Λ , let $\Lambda_r = \Lambda_{r,r}$. Let $A_0, A_1, \dots, A_k$ and B be $p^2 \times p^2$ matrices defined by, for $n \geq n_0$ with the notation of Assumption 9,

$$A_{0,i,j} = \begin{cases} -1/2 & \text{for } i = j \in \mathcal{D}(p), \\ -\sum_{l=1}^k \{D(f_l)_{I_p(i)} - D(f_l)_{J_p(i)}\} D(f_l)_{I_p(i)} & \text{for } i = j \notin \mathcal{D}(p), \\ 0 & \text{otherwise,} \end{cases}$$

$$A_{l,i,j} = \begin{cases} D(f_l)_{I_p(i)} - D(f_l)_{J_p(i)} & \text{for } i = j \notin \mathcal{D}(p), \\ 0 & \text{otherwise,} \end{cases} \quad \text{for } l = 1, \dots, k,$$

and

$$B_{i,j} = \begin{cases} 1 & \text{for } i = j \in \mathcal{D}(p), \\ [\sum_{l=1}^k \{D(f_l)_{I_p(i)} - D(f_l)_{J_p(i)}\}^2]^{-1} & \text{for } i = j \notin \mathcal{D}(p), \\ 0 & \text{otherwise.} \end{cases}$$

Let G be the $p^2 \times (k+1)p^2$ matrix defined by $G = B(A_0, A_1, \dots, A_k)$, for $n \geq n_0$. Let $M_{\Omega^{-1}}$ be as in Appendix A.3. Let $\tilde{V}(f_1, \dots, f_k)$ be the $(k+1)p^2 \times (k+1)p^2$ matrix composed of $(k+1)^2$ blocks of size $p^2 \times p^2$ with block $(i+1), (j+1)$ defined similarly as $\Sigma(f_i, f_j)$ in Appendix A.2, but with R replaced by R_z . Then, for $n \geq n_0$, F_k is defined as

$$F_k = M_{\Omega^{-1}} G \tilde{V}(f_1, \dots, f_k) G^T M_{\Omega^{-1}}^T.$$

A.5. Map for data application

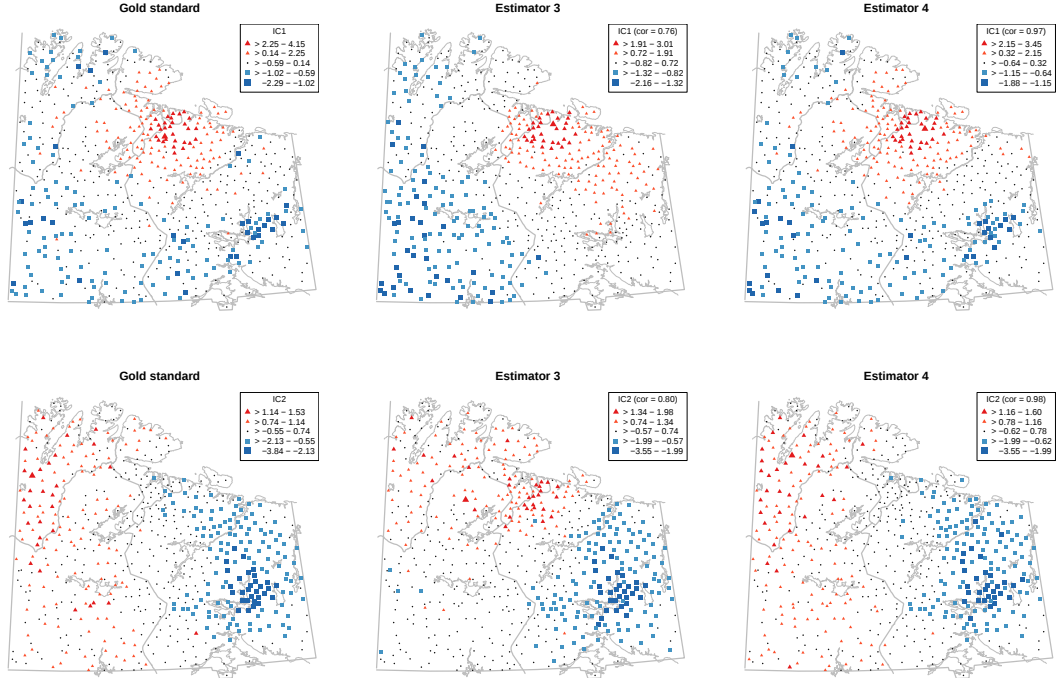


Fig. A.1: The first two independent components from the gold standard and estimators 3 and 4.

B. PROOFS

B.1. Introduction

We first prove Proposition 1 in Section B.2. Then Section B.3 provides general results for the proofs of Propositions 2, 3, 4, 6 and 7. Section B.4 provides the proofs of Propositions 2, 3 and 4. Section B.5 provides the proof of Proposition 5. Section B.6 provides the proofs of Propositions 6 and 7.

Propositions 2, 3, 4, 6 and 7 correspond to Proposition B.2, B.3, B.4, B.6 and B.7, respectively.

B.2. Proof of Proposition 1

We let $D(f) = \Omega^{-1} M(f) \Omega^{-T}$ for $f : \mathbb{R}^d \rightarrow \mathbb{R}$. We restate Proposition 1 and prove it.

PROPOSITION B.1. *The unmixing problem given by f is identifiable if and only if the diagonal elements of $\Omega^{-1}M(f)\Omega^{-\top}$ are distinct.*

Proof of Proposition B.1 (Proposition 1). We have that $\Gamma(f)$ is an unmixing functional if and only if

$$\Gamma(f)\Omega D(f_0)\Omega^\top \Gamma(f)^\top = I_p \quad \text{and} \quad \Gamma(f)\Omega D(f)\Omega^\top \Gamma(f)^\top = \Lambda(f).$$

Hence the matrix $\Gamma(f)\Omega$ is orthogonal, and its rows provide the eigenvectors of the diagonal matrix $D(f)$. If the diagonal elements of $D(f)$ are distinct, then the set of one-dimensional eigenspaces of $D(f)$ is unique and thus $\Gamma(f)\Omega = PS$ where P is a permutation matrix and S is a diagonal matrix with diagonal elements equal to -1 or 1 . If there are two diagonal elements of $D(f)$ that are equal, say the first and the second without loss of generality, then consider the $p \times p$ block diagonal matrix Q with first 2×2 block equal to $\{(2^{-1/2}, 2^{-1/2})^\top, (2^{-1/2}, -2^{-1/2})^\top\}$ and second $(p-2) \times (p-2)$ block equal to I_{p-2} . Then there is an unmixing functional $\Gamma(f)$ such that $\Gamma(f)\Omega = Q$. In this case, $\Gamma(f) = Q\Omega^{-1}$, which is not of the form $PS\Omega^{-1}$. $\square$

B.3. General results

Recall that $d \in \mathbb{N}$ and $p \in \mathbb{N}$ are fixed. $Z_1, \dots, Z_p$ are p independent stationary Gaussian processes on $\mathbb{R}^d$ with zero mean functions, unit variances and covariance functions $K_1, \dots, K_p$. We have $Z = (Z_1, \dots, Z_p)^\top$ and $X = (X_1, \dots, X_p)^\top = \Omega Z$ with Ω a fixed invertible $p \times p$ matrix.

Let $s_1, \dots, s_n$ be the n observation points in $\mathcal{S}^d$ and let f be a kernel function from $\mathbb{R}^d$ into $\mathbb{R}$. We recall

$$\widehat{M}(f) = n^{-1} \sum_{i=1}^n \sum_{j=1}^n f(s_i - s_j) X(s_i) X(s_j)^\top.$$

and

$$M(f) = n^{-1} \sum_{i=1}^n \sum_{j=1}^n f(s_i - s_j) \Omega D(s_i, s_j) \Omega^\top$$

where $D(s_i, s_j)$ is the $p \times p$ diagonal matrix defined by

$$D(s_i, s_j)_{k,k} = K_k(s_i - s_j).$$

Let $|x| = \max_{i=1, \dots, m} |x_i|$ be the sup norm for $x \in \mathbb{R}^m$. Since this norm is equivalent to the Euclidean norm, and since we work under Assumptions 5 to 7, we can assume without loss of generality that the following conditions hold.

Condition A1. With $\Delta > 0$ defined in Assumption 5, for all $n \in \mathbb{N}$ and for all $a \neq b$, $a, b \in \{1, \dots, n\}$, we have $|s_a - s_b| \geq \Delta$.

Condition A2. With $A < +\infty$ and $\alpha > 0$ defined in Assumptions 6 and 7, for all $s \in \mathbb{R}^d$ and for all $k = 1, \dots, p$, we have

$$|K_k(s)| \leq \frac{A}{1 + |s|^{d+\alpha}}.$$

Condition A3. With $A < +\infty$ and $\alpha > 0$ defined in Assumptions 6 and 7, for all $s \in \mathbb{R}^d$, we have

$$|f(s)| \leq \frac{A}{1 + |s|^{d+\alpha}}.$$

For a matrix M , denote by $M_{i,j}$ the element from the i th row and the j th column of M . For a vector V_n or a matrix M_n , denote by $(V_n)_i$ the i th element of V_n and by $(M_n)_{i,j}$ the element from the i th row and the j th column of M_n . The singular values of a $n \times n$ matrix M are denoted by $\rho_1(M) \geq \dots \geq \rho_n(M) \geq 0$ and, in the case when M is symmetric, the eigenvalues are denoted by $\lambda_1(M) \geq \dots \geq \lambda_n(M)$. The spectral norm is given by $\rho_1(M)$ and $\|M\|_F^2 = \sum_{i,j} (M_{i,j})^2$ denotes the Frobenius norm. For a sequence

of random variables X_n , we write $X_n = o_p(1)$ when X_n converges to 0 in probability as $n \rightarrow \infty$ and we write $X_n = O_p(1)$ when X_n is bounded in probability as $n \rightarrow \infty$. Let $e_i(k)$ be the i th base column vector of $\mathbb{R}^k$. Let y be the $np \times 1$ vector defined by $y_{(i-1)p+j} = X_j(s_i)$, for $i = 1, \dots, n, j = 1, \dots, p$.

LEMMA B.1. *Under Conditions A1 and A2, there exists a finite constant $C < +\infty$ so that for all $n \in \mathbb{N}$,*

$$\lambda_1\{\text{cov}(y)\} \leq C.$$

Proof. Let ℓ_a^T denote the a th row of Ω . We have

$$\begin{aligned} |\text{cov}\{X_a(s_i), X_b(s_j)\}| &= |\text{cov}\{\ell_a^T Z(s_i), \ell_b^T Z(s_j)\}| \\ &= |\ell_a^T D(s_i, s_j) \ell_b| \\ &\leq \sum_{k=1}^p |\Omega_{a,k}| |\Omega_{b,k}| \frac{A}{1 + |s_i - s_j|^{d+\alpha}} \\ &\leq p \max_{k,l=1,\dots,p} (\Omega_{k,l})^2 \frac{A}{1 + |s_i - s_j|^{d+\alpha}} \end{aligned} \quad (\text{A1})$$

from Condition A2. Note also that $\lambda_1\{\text{cov}(y)\} = \lambda_1\{\text{cov}(\tilde{y})\}$ where $\tilde{y}$ is the $np \times 1$ vector defined by $\tilde{y}_{(j-1)n+i} = X_j(s_i)$ for $i = 1, \dots, n, j = 1, \dots, p$. Hence, the lemma is a direct consequence of Lemma 6 in Furrer et al. (2016). $\square$

The next theorem provides a general multivariate central limit theorem for quadratic forms of Gaussian vectors. It extends standard central limit theorems in spatial statistics, see, e.g., Bachoc (2014) or Istas & Lang (1997), by allowing cases where the sequence of covariance matrices is non-converging or asymptotically singular. The full proof is given for self-consistency, although some of the arguments have appeared previously.

THEOREM B.1. *Let (y_n) be a sequence of n -dimensional centered Gaussian vectors. Let R_n be the covariance matrix of y_n . Assume that for all n , $\lambda_1(R_n) \leq A$ where A is a fixed finite constant. Let $k \in \mathbb{N}$ be fixed and let $(T_{1,n}), \dots, (T_{k,n})$ be k sequences of deterministic $n \times n$ symmetric matrices. Assume that for $i = 1, \dots, k$, for $n \in \mathbb{N}$, $\rho_1(T_{i,n}) \leq A$. Let Σ_n be the $k \times k$ matrix defined for $1 \leq i, j \leq k$, by*

$$(\Sigma_n)_{i,j} = 2n^{-1} \text{tr}(R_n T_{i,n} R_n T_{j,n}).$$

Let r_n be the k -dimensional vector defined for $i = 1, \dots, k$, by

$$(r_n)_i = \text{tr}(n^{-1} R_n T_{i,n}).$$

Let V_n be the $k \times 1$ vector defined for $1 \leq i \leq k$, by

$$(V_n)_i = n^{-1} y_n^T T_{i,n} y_n.$$

Let Q_n be the probability measure of $n^{1/2}(V_n - r_n)$ on $\mathbb{R}^k$. Let $\mathcal{N}(0, \Sigma_n)$ be the Gaussian distribution on $\mathbb{R}^k$ with mean vector 0 and covariance matrix Σ_n . Let d_w denote a metric generating the topology of weak convergence on the set of Borel probability measures on $\mathbb{R}^k$; for specific examples see the discussion in Dudley (2002) p. 393. Then we have, for $n \rightarrow \infty$,

$$d_w\{Q_n, \mathcal{N}(0, \Sigma_n)\} \rightarrow 0.$$

Proof. Assume that $d_w\{Q_n, \mathcal{N}(0, \Sigma_n)\} \not\rightarrow 0$ when $n \rightarrow \infty$. Then there exists $\epsilon > 0$ fixed and a subsequence n_m so that $d_w\{Q_{n_m}, \mathcal{N}(0, \Sigma_{n_m})\} \geq \epsilon$. Let $a_1, \dots, a_k \in \mathbb{R}$ be fixed. Let $S_{n_m} = R_{n_m}^{1/2} (\sum_{i=1}^k a_i T_{i,n_m}) R_{n_m}^{1/2}$. We have

$$\sum_{i,j=1}^k a_i a_j (\Sigma_{n_m})_{i,j} = 2 \text{tr}(S_{n_m}^2) / n_m.$$

Hence, we see that Σ_{n_m} is a non-negative matrix, and, from the assumptions on (R_{n_m}) and (T_{i,n_m}) , that $(\Sigma_{n_m})_{i,i} \leq 2A^4$. Also, $|(r_n)_i| = |n^{-1}\text{tr}(R_n^{1/2}T_{i,n}R_n^{1/2})| \leq A^2$. Hence, by compactity, and up to extracting a further subsequence, we can assume that $r_{n_m} \rightarrow r$ and $\Sigma_{n_m} \rightarrow \Sigma$ when $n_m \rightarrow \infty$. One can show simply that $d_w\{\mathcal{N}(0, \Sigma_{n_m}), \mathcal{N}(0, \Sigma)\} \rightarrow 0$ when $n_m \rightarrow \infty$. Hence, when $n_m \rightarrow \infty$,

$$\limsup d_w\{Q_{n_m}, \mathcal{N}(0, \Sigma)\} \geq \epsilon. \quad (\text{A2})$$

Let us prove (A2). We have

$$\begin{aligned} & |\limsup d_w\{Q_{n_m}, \mathcal{N}(0, \Sigma)\} - \limsup d_w\{Q_{n_m}, \mathcal{N}(0, \Sigma_{n_m})\}| \\ & \leq \limsup |d_w\{Q_{n_m}, \mathcal{N}(0, \Sigma)\} - d_w\{Q_{n_m}, \mathcal{N}(0, \Sigma_{n_m})\}| \\ & \leq \limsup d_w\{\mathcal{N}(0, \Sigma_{n_m}), \mathcal{N}(0, \Sigma)\} \\ & = 0, \end{aligned}$$

where we have applied the triangle inequality for the metric d_w in the last inequality above. Hence $\limsup d_w\{Q_{n_m}, \mathcal{N}(0, \Sigma)\} = \limsup d_w\{Q_{n_m}, \mathcal{N}(0, \Sigma_{n_m})\} \geq \epsilon$. Thus (A2) is proved.

We remark that the matrix $S_{n_m} = R_{n_m}^{1/2}(\sum_{i=1}^k a_i T_{i,n_m})R_{n_m}^{1/2}$ is symmetric, because $T_{1,n_m}, \dots, T_{k,n_m}$ are assumed to be symmetric in the theorem. Hence, S_{n_m} can be diagonalized and there exist a matrix P_{n_m} such that $P_{n_m} P_{n_m}^\top = I_{n_m}$ and a diagonal matrix D_{n_m} such that $S_{n_m} = P_{n_m} D_{n_m} P_{n_m}^\top$. Let also $z_{n_m} = R_{n_m}^{-1/2} y_{n_m}$. Observe that z_{n_m} follows the $\mathcal{N}(0, I_{n_m})$ distribution. We have

$$\begin{aligned} \sum_{i=1}^k a_i (V_n)_i &= n_m^{-1} y_{n_m}^\top \left(\sum_{i=1}^k a_i T_{i,n_m} \right) y_{n_m} \\ &= n_m^{-1} z_{n_m}^\top S_{n_m} z_{n_m} \\ &= n_m^{-1} \sum_{a=1}^{n_m} (\xi_{n_m})_a^2 \lambda_a(S_{n_m}), \end{aligned}$$

where ξ_{n_m} follows the $\mathcal{N}(0, I_{n_m})$ distribution. Hence letting

$$W_{n_m} = n_m^{1/2} \left\{ \sum_{i=1}^k a_i (V_{n_m})_i - \sum_{i=1}^k a_i (r_{n_m})_i \right\},$$

we have

$$W_{n_m} = n_m^{-1/2} \sum_{l=1}^{n_m} \{(\xi_{n_m})_l^2 - 1\} \lambda_l(S_{n_m}).$$

If $\sum_{i,j=1}^k a_i a_j \Sigma_{i,j} = 0$, then $\sum_{i,j=1}^k a_i a_j (\Sigma_{n_m})_{i,j} \rightarrow 0$ when $n_m \rightarrow \infty$. Hence, $2n_m^{-1} \text{tr}(S_{n_m}^2) \rightarrow 0$ and so $\text{var}(W_{n_m}) \rightarrow 0$. Hence $W_{n_m} \rightarrow \mathcal{N}(0, 0) = \mathcal{N}(0, \sum_{i,j=1}^k a_i a_j \Sigma_{i,j})$ in distribution when $n_m \rightarrow \infty$.

Now, if $\sum_{i,j=1}^k a_i a_j \Sigma_{i,j} > 0$, one can show from the Lindeberg-Feller central limit theorem that when $n_m \rightarrow \infty$, $W_{n_m} \rightarrow \mathcal{N}(0, \sum_{i,j=1}^k a_i a_j \Sigma_{i,j})$ in distribution, see also Lemma 2 in Istas & Lang (1997).

Hence, since both of the above-considered convergences in distribution hold for any $a_1, \dots, a_k$, we have, by Cramér-Wold theorem, that when $n_m \rightarrow \infty$, $n_m^{1/2}(V_{n_m} - r_{n_m}) \rightarrow \mathcal{N}(0, \Sigma)$ in distribution. This is in contradiction with (A2). Hence when $n \rightarrow \infty$

$$d_w\{Q_n, \mathcal{N}(0, \Sigma_n)\} \rightarrow 0.$$

□

B.4. Asymptotics when diagonalising two matrices

The next proposition gives the consistency of $\widehat{M}(f)$.

PROPOSITION B.2. *Let Conditions A1 and A2 hold and let $f : \mathbb{R}^d \rightarrow \mathbb{R}$ satisfy Condition A3. Then as $n \rightarrow \infty$, $\widehat{M}(f) - M(f) \rightarrow 0$ in probability.*

Proof. Clearly $E\{\widehat{M}(f)\} = M(f)$. Let $k, l \in \{1, \dots, p\}$ be fixed. In order to prove the proposition, it is sufficient to show that when $n \rightarrow \infty$, $\text{var}\{\widehat{M}(f)_{k,l}\} \rightarrow 0$.

We have,

$$\begin{aligned}\widehat{M}(f)_{k,l} &= n^{-1} \sum_{i=1}^n \sum_{j=1}^n f(s_i - s_j) e_k(p)^T X(s_i) X(s_j)^T e_l(p) \\ &= n^{-1} \sum_{i=1}^n \sum_{j=1}^n f(s_i - s_j) X(s_i)^T e_k(p) e_l(p)^T X(s_j) \\ &= n^{-1} \sum_{i=1}^n \sum_{j=1}^n f(s_i - s_j) X(s_i)^T [(1/2)\{e_k(p) e_l(p)^T + e_l(p) e_k(p)^T\}] X(s_j).\end{aligned}$$

Let $T_{k,l}(f)$ be the $np \times np$ matrix, that we see as a block matrix composed of n^2 blocks of sizes p^2 , and with block i, j equal to $f(s_i - s_j)(1/2)\{e_k(p) e_l(p)^T + e_l(p) e_k(p)^T\}$. We remark that $T_{k,l}(f)$ is symmetric. With this notation,

$$\widehat{M}(f)_{k,l} = n^{-1} y^T T_{k,l}(f) y.$$

The largest singular value of $T_{k,l}(f)$ is bounded as $n \rightarrow \infty$. Indeed, from Gershgorin's circle theorem, $\rho_1\{T_{k,l}(f)\}$ is no larger than $\max_{i=1, \dots, np} \sum_{j=1}^{np} |T_{k,l}(f)_{i,j}|$. This maximum is no larger than $\max_{i=1, \dots, n} \sum_{j=1}^n |f(s_i - s_j)|$. This last quantity is bounded as $n \rightarrow \infty$ from Condition A3 and from Lemma 4 in Furrer et al. (2016).

Hence $\rho_1\{T_{k,l}(f)\}$ is bounded by a constant $B < +\infty$. Thus, using Theorem 3.2d.3 in Mathai & Provost (1992) and the fact that $T_{k,l}(f)$ is symmetric,

$$\begin{aligned}\text{var}\{\widehat{M}(f)_{k,l}\} &= n^{-2} \text{var}\{y^T T_{k,l}(f) y\} \\ &= 2n^{-2} \text{tr}\{\text{cov}(y) T_{k,l}(f) \text{cov}(y) T_{k,l}(f)\} \\ &\leq 2pn^{-1} B^2 C^2,\end{aligned}$$

with $\lambda_1\{\text{cov}(y)\} \leq C$ from Lemma B.1. □

The next proposition is a corollary of Theorem B.1 and gives the asymptotic normality of $\widehat{M}(f)$.

PROPOSITION B.3. *Let, for $k, l = 1, \dots, p$ and $f : \mathbb{R}^d \rightarrow \mathbb{R}$, $T_{k,l}(f)$ be defined as in the proof of Proposition B.2. Let $R = \text{cov}(y)$ and let $\Sigma(f)$ be the $p^2 \times p^2$ matrix defined by, for $i = (s-1)p + t$ and $j = (u-1)p + v$, with $s, t, u, v \in \{1, \dots, p\}$,*

$$\Sigma(f)_{i,j} = 2n^{-1} \text{tr}\{RT(f)_{s,t} RT(f)_{u,v}\}.$$

Define, for $g : \mathbb{R}^d \rightarrow \mathbb{R}$, $\Sigma(f, g)$ as the $p^2 \times p^2$ matrix defined for $i = (s-1)p + t$ and $j = (u-1)p + v$, with $s, t, u, v \in \{1, \dots, p\}$ by

$$\Sigma(f, g)_{i,j} = 2n^{-1} \text{tr}\{RT(f)_{s,t} RT(g)_{u,v}\}.$$

Let

$$V(f, g) = \begin{pmatrix} \Sigma(f) & \Sigma(f, g) \\ \Sigma(g, f) & \Sigma(g) \end{pmatrix}.$$

Then, $V(f, g)$ is symmetric non-negative definite.

Assume that Conditions A1 and A2 hold. Let $f_1, f_2 : \mathbb{R}^d \rightarrow \mathbb{R}$ satisfy Condition A3. Let for $r = 1, 2$, $W(f_r)$ be the vector of size $p^2 \times 1$, defined for $i = (a-1)p + b$, $a, b \in \{1, \dots, p\}$, by $W(f_r)_i = n^{1/2}\{\widehat{M}(f_r)_{a,b} - M(f_r)_{a,b}\}$.

Let Q_n be the distribution of $\{W(f_1)^\top, W(f_2)^\top\}^\top$. Then as $n \rightarrow \infty$

$$d_w[Q_n, \mathcal{N}\{0, V(f_1, f_2)\}] \rightarrow 0. \quad (\text{A3})$$

Furthermore $\lambda_1\{V(f_1, f_2)\}$ is bounded as $n \rightarrow \infty$.

Proposition 3 is a direct corollary of Proposition B.3 with $f_1 = f$ and $f_2 = f_0$. Moreover, Proposition B.3 gives details concerning the matrix $V(f_1, f_2)$.

Proof. Let $a, b \in \{1, \dots, p\}$ and $f : \mathbb{R}^d \rightarrow \mathbb{R}$. We have seen in the proof of Proposition B.2 that

$$\widehat{M}(f)_{a,b} = n^{-1} y^\top T_{a,b}(f) y.$$

Hence

$$E(\widehat{M}(f)_{a,b}) = M(f)_{a,b} = n^{-1} \text{tr}\{RT_{a,b}(f)\}.$$

Let us first prove that $V(f, g)$ is symmetric non-negative definite. Let $W(f)$ be defined as $W(f_1)$, but with f_1 replaced by f . Let $W(g)$ be defined similarly with f_1 replaced by g . For $a_1, \dots, a_{p^2}, b_1, \dots, b_{p^2} \in \mathbb{R}$, we have

$$\begin{aligned} \text{var} \left\{ \sum_{i=1}^{p^2} a_i W(f)_i + \sum_{i=1}^{p^2} b_i W(g)_i \right\} &= \sum_{i,j=1}^{p^2} a_i a_j \text{cov}\{W(f)_i, W(f)_j\} + \sum_{i,j=1}^{p^2} b_i b_j \text{cov}\{W(g)_i, W(g)_j\} \\ &\quad + 2 \sum_{i,j=1}^{p^2} a_i b_j \text{cov}\{W(f)_i, W(g)_j\}. \end{aligned}$$

Now for $i, j = 1, \dots, p^2$, let $a, b, u, v \in \{1, \dots, p\}$ such that $i = (a-1)p + b$ and $j = (u-1)p + v$. We have, using Theorem 3.2d.3 in Mathai & Provost (1992) and the fact that $T_{a,b}(f)$ and $T_{u,v}(f)$ are symmetric,

$$\begin{aligned} \text{cov}\{W(f)_i, W(f)_j\} &= n \text{cov}\{\widehat{M}(f)_{a,b}, \widehat{M}(f)_{u,v}\} \\ &= n^{-1} \text{cov}\{y^\top T_{a,b}(f) y, y^\top T_{u,v}(f) y\} \\ &= 2n^{-1} \text{tr}\{RT_{a,b}(f) RT_{u,v}(f)\} \\ &= \Sigma(f)_{i,j}. \end{aligned}$$

We show similarly that

$$\text{cov}\{W(g)_i, W(g)_j\} = \Sigma(g)_{i,j}$$

and that

$$\text{cov}\{W(f)_i, W(g)_j\} = \text{cov}\{W(g)_j, W(f)_i\} = \Sigma(f, g)_{i,j} = \Sigma(g, f)_{j,i}.$$

Hence

$$\begin{aligned} \text{var} \left\{ \sum_{i=1}^{p^2} a_i W(f)_i + \sum_{i=1}^{p^2} b_i W(g)_i \right\} &= \sum_{i,j=1}^{p^2} a_i a_j \Sigma(f)_{i,j} + \sum_{i,j=1}^{p^2} b_i b_j \Sigma(g)_{i,j} \\ &\quad + \sum_{i,j=1}^{p^2} a_i b_j \Sigma(f, g)_{i,j} + \sum_{i,j=1}^{p^2} a_i b_j \Sigma(g, f)_{j,i} \\ &= (a_1, \dots, a_{p^2}, b_1, \dots, b_{p^2}) \begin{pmatrix} \Sigma(f) & \Sigma(f, g) \\ \Sigma(g, f) & \Sigma(g) \end{pmatrix} \begin{pmatrix} a_1 \\ \vdots \\ a_{p^2} \\ b_1 \\ \vdots \\ b_{p^2} \end{pmatrix}. \end{aligned}$$

Hence, since a square matrix is uniquely defined by its corresponding quadratic forms, it follows that $V(f, g)$ is the covariance matrix of the vector

$$(W(f)_1, \dots, W(f)_{p^2}, W(g)_1, \dots, W(g)_{p^2})^\top.$$

Hence, $V(f, g)$ is symmetric non-negative definite.

Let us now prove (A3). We have, from Lemma B.1 and the proof of Proposition B.2, that $\lambda_1(R)$ and $\rho_1\{T_{a,b}(f_r)\}$ are bounded as $n \rightarrow \infty$ for $r = 1, 2$. Hence, (A3) is a consequence of Theorem B.1. Finally, $\lambda_1\{V(f_1, f_2)\}$ is bounded as $n \rightarrow \infty$ because each component of $V(f_1, f_2)$ is bounded as $n \rightarrow \infty$. $\square$

Our objective is now to prove Proposition 4, which is a central limit theorem for $\widehat{\Gamma}(f) - \Omega^{-1}$ and $\widehat{\Lambda}(f) - \Lambda(f)$.

There is an equivariance property in Definition 1 that we will exploit. More precisely, let $D_0 = D(0, 0) = I_p$ and let

$$\widehat{D}_0 = n^{-1} \sum_{i=1}^n Z(s_i) Z(s_i)^\top.$$

For $f : \mathbb{R}^d \rightarrow \mathbb{R}$, let

$$D(f) = n^{-1} \sum_{i=1}^n \sum_{j=1}^n f(s_i - s_j) D(s_i, s_j)$$

and

$$\widehat{D}(f) = n^{-1} \sum_{i=1}^n \sum_{j=1}^n f(s_i - s_j) Z(s_i) Z(s_j)^\top.$$

Let $\Gamma\{\widehat{D}_0, \widehat{D}(f)\}$ and $\Lambda\{\widehat{D}_0, \widehat{D}(f)\}$ satisfy the following modification of Definition 1:

$$\Gamma\{\widehat{D}_0, \widehat{D}(f)\} \widehat{D}_0 \Gamma\{\widehat{D}_0, \widehat{D}(f)\}^\top = I_p \quad \text{and} \quad \Gamma\{\widehat{D}_0, \widehat{D}(f)\} \widehat{D}(f) \Gamma\{\widehat{D}_0, \widehat{D}(f)\}^\top = \Lambda\{\widehat{D}_0, \widehat{D}(f)\}, \quad (\text{A4})$$

where $\Lambda\{\widehat{D}_0, \widehat{D}(f)\}$ is a diagonal matrix with diagonal elements in decreasing order. Then, we can show that

$$\Gamma\{\widehat{M}_0, \widehat{M}(f)\} = \Gamma\{\widehat{D}_0, \widehat{D}(f)\} \Omega^{-1} \quad \text{and} \quad \Lambda\{\widehat{M}_0, \widehat{M}(f)\} = \Lambda\{\widehat{D}_0, \widehat{D}(f)\} \quad (\text{A5})$$

satisfy Definition 1. The above display is the equivariance property that we will exploit. That is, we will first show a central limit theorem for $\Gamma\{\widehat{D}_0, \widehat{D}(f)\} - I_p$ and $\Lambda\{\widehat{D}_0, \widehat{D}(f)\} - \Lambda(f)$ in Lemma B.4. Then,

we will use the equivariance property (A5) to obtain, directly, a central limit theorem for $\widehat{\Gamma}(f) - \Omega^{-1}$ and $\widehat{\Lambda}(f) - \Lambda(f)$ in Proposition B.4.

In the next lemma, we first show a central limit theorem for $\widehat{D}_0 - I_p$ and $\widehat{D}(f) - D(f)$. Recall that $\text{vect}(M) = (l_1^T, \dots, l_k^T)^T$ where $l_1^T, \dots, l_k^T$ are the k rows of a matrix M . Recall also the notation $f_0(x) = I(x = 0)$.

LEMMA B.2. *Let Conditions A1 and A2 hold. Let $f : \mathbb{R}^d \rightarrow \mathbb{R}$ satisfy Condition A3. Let*

$$Y_n = \begin{pmatrix} n^{1/2} \text{vect}(\widehat{D}_0 - D_0) \\ n^{1/2} \text{vect}\{\widehat{D}(f) - D(f)\} \end{pmatrix}.$$

Let $\tilde{V}(f)$ be as $V(f_0, f)$ in Proposition B.3 but where R is replaced by $\text{cov}(z)$ where z is the $np \times 1$ vector defined for $i = 1, \dots, n$, $j = 1, \dots, p$, by $z_{(i-1)p+j} = Z_j(s_i)$. Let $Q_{st,n}$ be the distribution of Y_n . Then we have

$$d_w[Q_{st,n}, \mathcal{N}\{0, \tilde{V}(f)\}] \rightarrow 0.$$

Furthermore, $\lambda_1(\tilde{V}(f))$ is bounded as $n \rightarrow \infty$.

Proof. The proof is identical to the proof of Proposition B.3. We remark that, with the notation of the proof of Proposition B.3, $W(f_r) = n^{1/2} \text{vect}\{\widehat{M}(f_r) - M(f_r)\}$, for $r = 1, 2$. $\square$

Now, we show, in Lemma B.3, that the transformation given by (A4), that defines $\Gamma\{\widehat{D}_0, \widehat{D}(f)\}$ and $\Lambda\{\widehat{D}_0, \widehat{D}(f)\}$ from $\widehat{D}_0$ and $\widehat{D}(f)$, is asymptotically linear, so to speak. This will allow us to transfer the central limit theorem for $\widehat{D}_0 - I_p$ and $\widehat{D}(f) - D(f)$ into a central limit theorem for $\Gamma\{\widehat{D}_0, \widehat{D}(f)\} - I_p$ and $\Lambda\{\widehat{D}_0, \widehat{D}(f)\} - \Lambda(f)$, for an appropriate choice of $\Gamma\{\widehat{D}_0, \widehat{D}(f)\}$. This argument is similar to the delta method in asymptotic statistics.

We will need the following notation. We let $\mathcal{D}(k) = \{1 + (i-1)(k+1); i = 1, \dots, k\}$. We remark that $\{\text{vect}(M)_i; i \in \mathcal{D}(k)\} = \{M_{i,i}; i = 1, \dots, k\}$ for a $k \times k$ matrix M . Let $\bar{\mathcal{D}}_k = \{1, \dots, k^2\} \setminus \mathcal{D}_k$. We remark that $\{\text{vect}(M)_i; i \in \bar{\mathcal{D}}(k)\} = \{M_{i,j}; i, j = 1, \dots, k, i \neq j\}$ for a $k \times k$ matrix M . For $a \in \{1, \dots, k^2\}$, let $I_k(a)$ and $J_k(a)$ be the unique $i, j \in \{1, \dots, k\}$ so that $a = k(i-1) + j$. For $i \in \{1, \dots, k\}$, let $d_k(i) = 1 + (i-1)(k+1)$ and note that $\{\text{vect}(M)_{d_k(i)}; i = 1, \dots, k\} = \{M_{i,i}; i = 1, \dots, k\}$ for a $k \times k$ matrix M . For a matrix M of size $k \times k$, recall that $\text{diag}(M) = (M_{1,1}, \dots, M_{k,k})^T$.

LEMMA B.3. *Let Conditions A1 and A2 hold. Let $f : \mathbb{R}^d \rightarrow \mathbb{R}$ satisfy Condition A3. Assume that Assumption 8 holds. Remark then that there exists $n_0 \in \mathbb{N}$ such that for $n \geq n_0$ the diagonal elements of $D(f)$ are strictly decreasing. Write these diagonal elements as $\lambda_1 > \dots > \lambda_p$. For $n \geq n_0$, let A be the $p^2 \times p^2$ matrix defined by*

$$A_{i,j} = \begin{cases} -1/2 & \text{for } i = j \in \mathcal{D}(p), \\ -\lambda_{I_p(i)} \{\lambda_{I_p(i)} - \lambda_{J_p(i)}\}^{-1} & \text{for } i = j \notin \mathcal{D}(p), \\ 0 & \text{otherwise.} \end{cases}$$

Let B be the $p^2 \times p^2$ matrix defined by

$$B_{i,j} = \begin{cases} \{\lambda_{I_p(i)} - \lambda_{J_p(i)}\}^{-1} & \text{for } i = j \notin \mathcal{D}(p), \\ 0 & \text{otherwise.} \end{cases}$$

Let C be the $p \times p^2$ matrix defined by

$$C_{i,j} = \begin{cases} -\lambda_i & \text{for } j = d_p(i), \\ 0 & \text{otherwise.} \end{cases}$$

Let D be the $p \times p^2$ matrix defined by

$$D_{i,j} = \begin{cases} 1 & \text{for } j = d_p(i), \\ 0 & \text{otherwise.} \end{cases}$$

Then, with probability going to one as $n \rightarrow \infty$, there exist $\Gamma\{\hat{D}_0, \hat{D}(f)\}$ and $\Lambda\{\hat{D}_0, \hat{D}(f)\}$ satisfying (A4). Furthermore, $\Gamma\{\hat{D}_0, \hat{D}(f)\}$ can be chosen so that as $n \rightarrow \infty$,

$$\begin{pmatrix} n^{1/2}(\text{vect}[\Gamma\{\hat{D}_0, \hat{D}(f)\} - I_p]) \\ n^{1/2}(\text{diag}[\Lambda\{\hat{D}_0, \hat{D}(f)\} - D(f)]) \end{pmatrix} = \begin{pmatrix} A & B \\ C & D \end{pmatrix} \begin{pmatrix} n^{1/2}\{\text{vect}(\hat{D}_0 - D_0)\} \\ n^{1/2}[\text{vect}\{\hat{D}(f) - D(f)\}] \end{pmatrix} + o_p(1).$$

Proof. Let us assume that $n \geq n_0$ throughout the proof. From Proposition B.2, with probability going to one, the eigenvalues of $D_0^{-1/2}\hat{D}(f)D_0^{-1/2}$ are distinct. In the rest of the proof, we set ourselves on the event when this is the case. Then, choose $\hat{\Gamma} = \Gamma\{\hat{D}_0, \hat{D}(f)\}$ and $\hat{\Lambda} = \Lambda\{\hat{D}_0, \hat{D}(f)\}$ satisfying (A4) and such that

$$\sum_{j=1}^p \hat{\Gamma}_{i,j} \geq 0 \quad (i = 1, \dots, p). \quad (\text{A6})$$

We remark that $\hat{\Gamma}$ and $\hat{\Lambda}$ indeed exist, since when $\Gamma\{\hat{D}_0, \hat{D}(f)\}$ and $\Lambda\{\hat{D}_0, \hat{D}(f)\}$ satisfy (A4), one can multiply each row of $\Gamma\{\hat{D}_0, \hat{D}(f)\}$ by 1 or -1 and still satisfy (A4).

Let

$$T_1 = \begin{pmatrix} n^{1/2}\text{vect}(\hat{\Gamma} - I_p) \\ n^{1/2}\text{diag}\{\hat{\Lambda} - D(f)\} \end{pmatrix}$$

and

$$T_2 = \begin{pmatrix} A & B \\ C & D \end{pmatrix} \begin{pmatrix} n^{1/2}\{\text{vect}(\hat{D}_0 - D_0)\} \\ n^{1/2}[\text{vect}\{\hat{D}(f) - D(f)\}] \end{pmatrix}.$$

Assume that $T_1 - T_2 \not\rightarrow 0$ in probability when $n \rightarrow \infty$. Then there exist $\epsilon > 0$ and a subsequence $n_m \rightarrow \infty$ so that along n_m

$$P(|T_1 - T_2| \geq \epsilon) \geq \epsilon. \quad (\text{A7})$$

One can show, as for the proof of Proposition B.2, that $\limsup \lambda_1\{D(f)\} < +\infty$ when $n \rightarrow \infty$. Hence, up to extracting a further subsequence, we can assume that when $n_m \rightarrow \infty$, $D(f) \rightarrow D_\infty(f)$, where $D_\infty(f)$ has distinct, decreasing, eigenvalues.

From Lemma 4.3 in Sun & Sun (2002), since $D_0^{-1/2}D_\infty(f)D_0^{-1/2} = D_\infty(f)$ is diagonal, there exists a sequence of random orthogonal matrices U_n such that $U_{n_m}\hat{D}_0^{-1/2}\hat{D}(f)\hat{D}_0^{-1/2}U_{n_m}^\top = \Lambda_{n_m}$ is diagonal and goes to $D_\infty(f)$ in probability and so that $U_{n_m} \rightarrow I_p$ in probability when $n_m \rightarrow \infty$. Hence, the pair $(U_{n_m}\hat{D}_0^{-1/2}, \Lambda_{n_m})$ satisfies (A4) with probability going to one. Furthermore, all the matrices $\Gamma\{\hat{D}_0, \hat{D}(f)\}$ for which $\Gamma\{\hat{D}_0, \hat{D}(f)\}$ and $\hat{\Lambda}$ satisfy (A4) satisfy $\Gamma\{\hat{D}_0, \hat{D}(f)\} = SU_{n_m}\hat{D}_0^{-1/2}$ where S is a diagonal matrix with diagonal elements equal to -1 or 1 . Hence with probability going to one, we must have $S = I_p$ for (A6) to be also satisfied. Hence, with probability going to one, $\hat{\Gamma} = U_{n_m}\hat{D}_0^{-1/2}$. Hence we have finally obtained $\hat{\Gamma} \rightarrow I_p$ and $|\hat{\Lambda} - D(f)| \rightarrow 0$ in probability when $n_m \rightarrow \infty$.

The rest of the proof is similar to those given in Ilmonen et al. (2010a) and Miettinen et al. (2012). By definition of $\hat{\Gamma}$ and $\hat{\Lambda}$, we have

$$\hat{\Gamma}\hat{D}_0\hat{\Gamma}^\top = I_p \quad \text{and} \quad \hat{\Gamma}\hat{D}(f)\hat{\Gamma}^\top = \hat{\Lambda}.$$

Hence

$$\begin{aligned} &(\widehat{\Gamma} - I_p)\widehat{D}_0\widehat{\Gamma}^\top + (\widehat{D}_0 - I_p)\widehat{\Gamma}^\top + (\widehat{\Gamma} - I_p)^\top = 0 \text{ and} \\ &(\widehat{\Gamma} - I_p)\widehat{D}(f)\widehat{\Gamma}^\top + \{\widehat{D}(f) - D(f)\}\widehat{\Gamma}^\top + D(f)(\widehat{\Gamma} - I_p)^\top = \widehat{\Lambda} - D(f). \end{aligned}$$

Also, from Lemma B.2, we have $n_m^{1/2}(\widehat{D}_0 - I_p) = O_p(1)$ and $n_m^{1/2}\{\widehat{D}(f) - D(f)\} = O_p(1)$. Thus, we get

$$\begin{aligned} n_m^{1/2}(\widehat{D}_0 - I_p) &= -n_m^{1/2}(\widehat{\Gamma} - I_p) - n_m^{1/2}(\widehat{\Gamma} - I_p)^\top + o_p(1) \text{ and} \\ n_m^{1/2}\{\widehat{D}(f) - D(f)\} &= -n_m^{1/2}(\widehat{\Gamma} - I_p)D(f) - n_m^{1/2}D(f)(\widehat{\Gamma} - I_p)^\top + n_m^{1/2}\{\widehat{\Lambda} - D(f)\} + o_p(1). \end{aligned}$$

This then yields

$$\begin{aligned} n_m^{1/2}(\widehat{\Gamma}_{ii} - 1) &= -\frac{1}{2}n_m^{1/2}(\widehat{D}_{0,i,i} - 1) + o_p(1) \\ (\lambda_i - \lambda_j)n_m^{1/2}\widehat{\Gamma}_{i,j} &= n_m^{1/2}\widehat{D}(f)_{i,j} - \lambda_i n_m^{1/2}\widehat{D}_{0,i,j} + o_p(1), \quad i \neq j, \text{ and} \\ n_m^{1/2}(\widehat{\Lambda}_{i,i} - \lambda_i) &= n_m^{1/2}\{\widehat{D}(f)_{i,i} - \lambda_i\} - \lambda_i n_m^{1/2}(\widehat{D}_{0,i,i} - 1) + o_p(1). \end{aligned}$$

This is in contradiction with (A7), by definition of A , B , C and D . Hence the proof is finished. $\square$

From Lemmas B.2 and B.3, we now obtain a central limit theorem for $\Gamma\{\widehat{D}_0, \widehat{D}(f)\} - I_p$ and $\Lambda\{\widehat{D}_0, \widehat{D}(f)\} - \Lambda(f)$. Note that $\Lambda(f) = D(f)$.

LEMMA B.4. *Assume the same conditions as in Lemma B.3 and let n_0 be defined as in Lemma B.3. Let, for $n \geq n_0$,*

$$G = \begin{pmatrix} A & B \\ C & D \end{pmatrix},$$

from Lemma B.3. For $\Gamma\{\widehat{D}_0, \widehat{D}(f)\}$ and $\Lambda\{\widehat{D}_0, \widehat{D}(f)\}$ satisfying (A4), let

$$X_n = n^{1/2} \begin{pmatrix} \text{vect}[\Gamma\{\widehat{D}_0, \widehat{D}(f)\} - I_p] \\ \text{diag}[\Lambda\{\widehat{M}_0, \widehat{M}(f)\} - D(f)] \end{pmatrix}.$$

Let Q_n be the distribution of X_n . Let $\tilde{V}(f)$ be defined as in Lemma B.2. Then, we can choose $\Gamma\{\widehat{D}_0, \widehat{D}(f)\}$ and $\Lambda\{\widehat{D}_0, \widehat{D}(f)\}$ satisfying (A4) such that when $n \rightarrow \infty$,

$$d_w\{Q_n, \mathcal{N}(0, G\tilde{V}(f)G^\top)\} \rightarrow 0.$$

Proof. The lemma is a direct consequence of Lemmas B.2 and B.3. The proof is carried out by contradiction, by taking subsequences along which the bounded sequences of matrices G and $\tilde{V}(f)$ converge, and by applying Slutsky's lemma. $\square$

We now use the equivariance property (A5) to conclude.

PROPOSITION B.4. *Assume the same conditions as in Lemma B.3. Let $\tilde{V}(f)$ and G be defined as in Lemmas B.2 and B.4. Let $M_{\Omega^{-1}}$ be the $p^2 \times p^2$ matrix defined by*

$$(M_{\Omega^{-1}})_{a,b} = \begin{cases} (\Omega^{-1})_{J_p(b), J_p(a)} & \text{if } I_p(a) = I_p(b), \\ 0 & \text{if } I_p(a) \neq I_p(b). \end{cases}$$

Let $\bar{M}_{\Omega^{-1}}$ be the matrix of size $(p^2 + p) \times (p^2 + p)$ defined by

$$\bar{M}_{\Omega^{-1}} = \begin{pmatrix} M_{\Omega^{-1}} & 0 \\ 0 & I_p \end{pmatrix}.$$

For $\Gamma\{\widehat{M}_0, \widehat{M}(f)\}$ and $\Lambda\{\widehat{M}_0, \widehat{M}(f)\}$ satisfying Definition 1, let

$$X_n = n^{1/2} \begin{pmatrix} \text{vect}[\Gamma\{\widehat{M}_0, \widehat{M}(f)\} - \Omega^{-1}] \\ \text{diag}[\Lambda\{\widehat{M}_0, \widehat{M}(f)\} - \Lambda(f)] \end{pmatrix}.$$

Let Q_n be the distribution of X_n . Let, for $n \geq n_0$ with n_0 as in Lemma B.3,

$$F = \bar{M}_{\Omega^{-1}} G \tilde{V}(f) G^T \bar{M}_{\Omega^{-1}}^T.$$

Then, we can choose $\Gamma\{\widehat{M}_0, \widehat{M}(f)\}, \Lambda\{\widehat{M}_0, \widehat{M}(f)\}$, satisfying Definition 1, so that when $n \rightarrow \infty$,

$$d_w\{Q_n, \mathcal{N}(0, F)\} \rightarrow 0.$$

Proof. The proof directly follows from Lemma B.4 and from (A5). Indeed, for $\Gamma\{\widehat{D}_0, \widehat{D}(f)\}$ and $\Lambda\{\widehat{D}_0, \widehat{D}(f)\}$ satisfying the central limit theorem in Lemma B.4, we can choose $\Gamma\{\widehat{M}_0, \widehat{M}(f)\}$ and $\Lambda\{\widehat{M}_0, \widehat{M}(f)\}$ satisfying Definition 1 and such that

$$\begin{pmatrix} \text{vect}[\Gamma\{\widehat{M}_0, \widehat{M}(f)\} - \Omega^{-1}] \\ \text{diag}[\Lambda\{\widehat{M}_0, \widehat{M}(f)\} - D(f)] \end{pmatrix} = \bar{M}_{\Omega^{-1}} \begin{pmatrix} \text{vect}[\Gamma\{\widehat{D}_0, \widehat{D}(f)\} - I_p] \\ \text{diag}[\Lambda\{\widehat{D}_0, \widehat{D}(f)\} - D(f)] \end{pmatrix}.$$

We also remark that $\Lambda(f) = D(f)$. $\square$

B.5. Proof of Proposition 5

We let $D(f) = \Omega^{-1} M(f) \Omega^{-T}$ for $f : \mathbb{R}^d \rightarrow \mathbb{R}$. We restate Proposition 5 and prove it.

PROPOSITION B.5. *The unmixing problem given by $f_1, \dots, f_k$ is identifiable if and only if for every pair $i \neq j$, $i, j = 1, \dots, p$, there exists $l = 1, \dots, k$ such that $\{\Omega^{-1} M(f_l) \Omega^{-T}\}_{i,i} \neq \{\Omega^{-1} M(f_l) \Omega^{-T}\}_{j,j}$.*

Proof of Proposition B.5 (Proposition 5). We have that Γ satisfies (6) if and only if

$$\Gamma \Omega \in \underset{\substack{L: LL^T = I_p \\ L \text{ has rows } l_1^T, \dots, l_p^T}}{\text{argmax}} \sum_{l=1}^k \sum_{j=1}^p \{l_j^T D(f_l) l_j\}^2. \quad (\text{A8})$$

If the condition of the proposition holds, then only orthogonal matrices of the form PS satisfy (A8), with the double sum being equal to

$$\sum_{l=1}^k \sum_{j=1}^p D(f_l)_{j,j}^2,$$

where P is a permutation matrix and S is a diagonal matrix with diagonal elements equal to -1 or 1 , see the end of the proof of Lemma B.5 below. Assume now that there exist $i \neq j$, $i, j = 1, \dots, p$, such that for $l = 1, \dots, k$, $\{\Omega^{-1} M(f_l) \Omega^{-T}\}_{i,i} = \{\Omega^{-1} M(f_l) \Omega^{-T}\}_{j,j}$. Without loss of generality, assume that $i = 1$ and $j = 2$. Then consider the $p \times p$ block diagonal matrix Q with first 2×2 block equal to $\{(2^{-1/2}, 2^{-1/2})^T, (2^{-1/2}, -2^{-1/2})^T\}$ and second $(p-2) \times (p-2)$ block equal to I_{p-2} . Then one can show that Q satisfies (A8). In this case, $\Gamma = Q\Omega^{-1}$ satisfies (6), and is not of the form $PS\Omega^{-1}$. $\square$

B.6. Asymptotics when diagonalising more than two matrices

LEMMA B.5. *Let Conditions A1 and A2 hold. Let $k \in \mathbb{N}$ be fixed. Let $f_1, \dots, f_k : \mathbb{R}^d \rightarrow \mathbb{R}$ satisfy Condition A3. Assume that Assumption 9 holds. Let $\widehat{\Gamma} = \widehat{\Gamma}\{\widehat{D}_0, \widehat{D}(f_1), \dots, \widehat{D}(f_k)\}$ be such that*

$$\widehat{\Gamma} \in \underset{\substack{\Gamma: \Gamma \widehat{D}_0 \Gamma^T = I_p \\ \Gamma \text{ has rows } \gamma_1^T, \dots, \gamma_p^T}}{\text{argmax}} \sum_{l=1}^k \sum_{j=1}^p \{\gamma_j^T \widehat{D}(f_l) \gamma_j\}^2. \quad (\text{A9})$$

Then we can choose $\widehat{\Gamma}$ so that $\widehat{\Gamma} \rightarrow I_p$ in probability when $n \rightarrow \infty$.

Proof. Let, for U a $p \times p$ orthogonal matrix with rows $u_1^T, \dots, u_p^T$,

$$\widehat{g}(U) = \sum_{l=1}^k \sum_{j=1}^p \{u_j^T \widehat{D}_0^{-1/2} \widehat{D}(f_l) \widehat{D}_0^{-1/2} u_j\}^2.$$

Let

$$E_0 = \{U \text{ orthogonal with rows } u_1^T, \dots, u_p^T; \\ \sum_{j=1}^p j(u_1)_j^2 \leq \dots \leq \sum_{j=1}^p j(u_p)_j^2 \text{ and for } i = 1, \dots, p, \sum_{j=1}^p U_{i,j} \geq 0\}.$$

We observe that any orthogonal matrix can be obtained from a matrix in E_0 , by row permutation and row multiplication by 1 or -1 . Hence, for any n , there exists $\widehat{U}$ so that $\widehat{U} \in \operatorname{argmax}_{U \in E_0} \widehat{g}(U)$ and $\widehat{U} \widehat{D}_0^{-1/2}$ satisfies (A9).

We now aim at showing that $\widehat{U} \rightarrow I_p$ in probability as $n \rightarrow \infty$, which will conclude the proof since $\widehat{D}_0 \rightarrow I_p$ in probability. Assume that this is not the case. Then, there exists $\epsilon > 0$ and a subsequence $(n_m)_{m \in \mathbb{N}}$ so that for all $m \in \mathbb{N}$ and along n_m

$$\operatorname{pr}(\|\widehat{U} - I_p\|_F \geq \epsilon) \geq \epsilon. \quad (\text{A10})$$

The matrices $D(f_1), \dots, D(f_l)$ are bounded (this can be shown as in Proposition B.2). Hence, by compactness, up to extracting a further subsequence, we have that (A10) holds along n_m and, as $m \rightarrow \infty$ and along n_m , $D(f_1) \rightarrow D_\infty(f_1), \dots, D(f_k) \rightarrow D_\infty(f_k)$.

We let

$$g_\infty(U) = \sum_{l=1}^k \sum_{j=1}^p \{u_j^T D_\infty(f_l) u_j\}^2.$$

We have, from Proposition B.2 and as observed in Miettinen et al. (2016), that, as $m \rightarrow \infty$ and along n_m ,

$$\sup_{U \in E_0} |\widehat{g}(U) - g_\infty(U)| \rightarrow 0$$

in probability as $m \rightarrow \infty$. Hence, using a standard M-estimator argument and because E_0 is compact, if the unique maximum of g_∞ on E_0 is I_p , we obtain that, as $m \rightarrow \infty$ and along n_m , $\widehat{U} \rightarrow I_p$ in probability. This is contradictory to (A10).

Hence, to conclude the proof, it suffices to show that the unique maximum of g_∞ on E_0 is I_p . We have

$$\begin{aligned} g_\infty(U) &= \sum_{l=1}^k \|U^T D_\infty(f_l) U\|_F^2 - \sum_{i \neq j} \{U^T D_\infty(f_l) U\}_{i,j}^2 \\ &\leq \sum_{l=1}^k \|U^T D_\infty(f_l) U\|_F^2 \\ &= \sum_{l=1}^k \|D_\infty(f_l)\|_F^2. \end{aligned} \quad (\text{A11})$$

Also,

$$g_\infty(I_p) = \sum_{l=1}^k \sum_{j=1}^p D_\infty(f_l)_{j,j}^2 = \sum_{l=1}^k \|D_\infty(f_l)\|_F^2.$$

We next show that the identity matrix I_p is the unique maximizer of g_∞ in E_0 . To see this, consider an arbitrary orthogonal matrix U which maximizes g_∞ . From (A11) we see that $U^T D_\infty(f_l) U$ is a diagonal matrix for all $l = 1, \dots, k$. Then, by its non-singularity, the matrix U must have a column with a non-zero first element. Call the first (from the left) such column of U by u . We show that all other elements of u must be zero. By the previous, u is an eigenvector of all $D_\infty(f_l)$ and we have,

$$D_\infty(f_l)u = \psi_l u, \quad \text{for all } l = 1, \dots, k,$$

for some eigenvalues $\psi_l \in \mathbb{R}$, $l = 1, \dots, k$. Assume then that u has a second non-zero element at some arbitrary position $q \neq 1$, meaning that both $u_1, u_q \neq 0$. Then we write

$$D_\infty(f_l)_{1,1}u_1 = \psi_l u_1 \quad \text{and} \quad D_\infty(f_l)_{q,q}u_q = \psi_l u_q, \quad \text{for all } l = 1, \dots, k,$$

which in turn implies that $D_\infty(f_l)_{1,1} = D_\infty(f_l)_{q,q}$ for all $l = 1, \dots, k$. By a continuity argument, this is a contradiction with Assumption 9. As the choice of q was arbitrary, the only non-zero element in u is the first. Repeating now the same reasoning for other elements besides the first, we observe that each column of the maximizer U must have a single non-zero element, and by its orthogonality we have $U = PD$ for some permutation matrix P and some diagonal matrix D with diagonal components in $\{-1, 1\}$. The only matrix of that form belonging to E_0 is I_p and thus, for all $U \in E_0$ with $U \neq I_p$, we have $g(U) < g(I_p)$. $\square$

LEMMA B.6. Assume the same setting and conditions as in Lemma B.5. For a diagonal matrix Λ , let $\Lambda_r = \Lambda_{r,r}$. Let $A_0, A_1, \dots, A_k$ and B be $p^2 \times p^2$ matrices defined by, for $n \geq n_0$ with the notation of Assumption 9,

$$A_{0,i,j} = \begin{cases} -1/2 & \text{for } i = j \in \mathcal{D}(p), \\ -\sum_{l=1}^k \{D(f_l)_{I_p(i)} - D(f_l)_{J_p(i)}\} D(f_l)_{I_p(i)} & \text{for } i = j \notin \mathcal{D}(p), \\ 0 & \text{otherwise,} \end{cases}$$

for $l = 1, \dots, k$,

$$A_{l,i,j} = \begin{cases} D(f_l)_{I_p(i)} - D(f_l)_{J_p(i)} & \text{for } i = j \notin \mathcal{D}(p), \\ 0 & \text{otherwise,} \end{cases}$$

and

$$B_{i,j} = \begin{cases} 1 & \text{for } i = j \in \mathcal{D}(p), \\ [\sum_{l=1}^k \{D(f_l)_{I_p(i)} - D(f_l)_{J_p(i)}\}^2]^{-1} & \text{for } i = j \notin \mathcal{D}(p), \\ 0 & \text{otherwise.} \end{cases}$$

Then, as $n \rightarrow \infty$, there exists $\hat{\Gamma}$ satisfying (A9) so that

$$n^{1/2} \text{vect}(\hat{\Gamma} - I_p) = B \begin{pmatrix} A_0 & A_1 & \dots & A_k \end{pmatrix} \begin{pmatrix} n^{1/2} \text{vect}(\hat{D}_0 - D_0) \\ n^{1/2} \text{vect}\{\hat{D}(f_1) - D(f_1)\} \\ \vdots \\ n^{1/2} \text{vect}\{\hat{D}(f_k) - D(f_k)\} \end{pmatrix}.$$

Proof. Assume that $n \geq n_0$ throughout the proof. Let $\hat{\Gamma}$ satisfy (A9) and $\hat{\Gamma} \rightarrow I_p$ in probability when $n \rightarrow \infty$ (the existence follows from Lemma B.5). The proof of the lemma follows the proofs of ii) in Theorem 2 of Miettinen et al. (2016) and Theorem 3 in Virta et al. (2018) and as such, we present below only some key steps.

From Lemma B.2, we have $n^{1/2}(\hat{D}_0 - I_p) = O_p(1)$ and $n^{1/2}\{\hat{D}(f_l) - D(f_l)\} = O_p(1)$, for all $l = 1, \dots, k$. By a continuity argument and our assumptions, we further have $D(f_1) \rightarrow D_\infty(f_1), \dots, D(f_k) \rightarrow D_\infty(f_k)$ such that the limit matrices satisfy: there exists a fixed $\delta > 0$ so that for every pair $i \neq j$, $i, j = 1, \dots, p$, there exists $l = 1, \dots, k$ such that $|D_\infty(f_l)_{i,i} - D_\infty(f_l)_{j,j}| \geq \delta$. The previous convergence holds up to extracting a subsequence. We omit this step in this proof for concision,

but see the proof of Lemma B.3. Finally, the rotation $\widehat{U}$ so that $\widehat{\Gamma} = \widehat{U}\widehat{D}_0^{-1/2}$ also satisfies $\widehat{U} \rightarrow I_p$ in probability.

Then, as in Virtà et al. (2018), the maximisation problem,

$$\underset{\substack{U: U U^T = I_p \\ U \text{ has rows } u_1^T, \dots, u_p^T}}{\operatorname{argmax}} \sum_{l=1}^k \sum_{j=1}^p \{u_j^T \widehat{D}_0^{-1/2} \widehat{D}(f_l) \widehat{D}_0^{-1/2} u_j\}^2,$$

yields the estimation equations $n^{1/2}\widehat{Y} = n^{1/2}\widehat{Y}^T$, where

$$n^{1/2}\widehat{Y} = n^{1/2} \sum_{l=1}^k \widehat{U} \widehat{R}(f_l) \widehat{U}^T \operatorname{Diag}\{\widehat{U} \widehat{R}(f_l) \widehat{U}^T\},$$

where we have used the shorthand $\widehat{R}(f_l) = \widehat{D}_0^{-1/2} \widehat{D}(f_l) \widehat{D}_0^{-1/2}$ and $\operatorname{Diag}(M)$ is equal to the square matrix M but with its off-diagonal elements set to zero. Linearizing the estimating equations asymptotically and vectorizing, we arrive at the following form,

$$(I_p - K) \sum_{l=1}^k \left([\operatorname{Diag}\{\widehat{U} D(f_l) \widehat{U}^T\} \widehat{U} D(f_l) \otimes I_p] + [\operatorname{Diag}\{\widehat{U} D(f_l) \widehat{U}^T\} \otimes D(f_l)] K \right) \cdot n^{1/2} \operatorname{vec}(\widehat{U} - I_p) = -(I_p - K) n^{1/2} \operatorname{vec}(\widehat{F}) + o_p(1), \quad (\text{A12})$$

where K is the $p^2 \times p^2$ commutation matrix satisfying $K^2 = I_p$, $n^{1/2}\widehat{F} = \sum_{l=1}^k n^{1/2}\{\widehat{R}(f_l) - D(f_l)\} D_\infty(f_l) = O_p(1)$ and $\operatorname{vec}(M) = (c_1^T, \dots, c_q^T)^T$ is the column vectorization where $c_1, \dots, c_q$ are the q columns of a matrix M . The orthogonality constraint can be similarly linearized to yield,

$$\{(\widehat{U} \otimes I_p) + K\} n^{1/2} \operatorname{vec}(\widehat{U} - I_p) = 0. \quad (\text{A13})$$

Summing (A12) and (A13), we obtain,

$$\widehat{A} n^{1/2} \operatorname{vec}(\widehat{U} - I_p) = -(I_p - K) n^{1/2} \operatorname{vec}(\widehat{F}) + o_p(1), \quad (\text{A14})$$

where $\widehat{A} \rightarrow A = (I_p - K)\{(\sum_{l=1}^k D_l^2 \otimes I_p) + (\sum_{l=1}^k D_l \otimes D_l)K\} + I_p + K$ in probability, where we use the notation $D_l = D_\infty(f_l)$, $l = 1, \dots, k$. Using the fact that $K(A \otimes B)K = B \otimes A$ for any conformable matrices A, B , we get the alternative form,

$$A = \left(\sum_{l=1}^k D_l^2 \otimes I_p - \sum_{l=1}^k D_l \otimes D_l + I_p \right) + \left(\sum_{l=1}^k D_l \otimes D_l - \sum_{l=1}^k I_p \otimes D_l^2 + I_p \right) K.$$

Continuing as in Virtà et al. (2018), each diagonal element of $\widehat{U}$ has a corresponding 1×1 diagonal block equal to 2 in A . Similarly, each pair of (a, b) th and (b, a) th off-diagonal elements in $\widehat{U}$ has a corresponding 2×2 sub-matrix in A of the form,

$$A_{ab} = \begin{pmatrix} 1 + \sum_{l=1}^k d_{la}^2 - \sum_{l=1}^k d_{la} d_{lb} & 1 - \sum_{l=1}^k d_{lb}^2 + \sum_{l=1}^k d_{la} d_{lb} \\ 1 - \sum_{l=1}^k d_{la}^2 + \sum_{l=1}^k d_{la} d_{lb} & 1 + \sum_{l=1}^k d_{lb}^2 - \sum_{l=1}^k d_{la} d_{lb} \end{pmatrix},$$

where d_{la} is the a th diagonal element of D_l . The inverse of the sub-matrix is

$$A_{ab}^{-1} = \left\{ 2 \sum_{l=1}^k (d_{la} - d_{lb})^2 \right\}^{-1} \begin{pmatrix} 1 + \sum_{l=1}^k d_{lb}^2 - \sum_{l=1}^k d_{la} d_{lb} & \sum_{l=1}^k d_{lb}^2 - \sum_{l=1}^k d_{la} d_{lb} - 1 \\ \sum_{l=1}^k d_{la}^2 - \sum_{l=1}^k d_{la} d_{lb} - 1 & 1 + \sum_{l=1}^k d_{la}^2 - \sum_{l=1}^k d_{la} d_{lb} \end{pmatrix},$$

showing that A is invertible as by our assumptions $\sum_{l=1}^k (d_{la} - d_{lb})^2 \neq 0$ for all distinct pairs $a, b = 1, \dots, p$. Thus, by Slutsky's theorem, we obtain from (A14) that,

$$n^{1/2} \operatorname{vec}(\widehat{U} - I_p) = -A^{-1} n^{1/2} \operatorname{vec}(\widehat{F} - \widehat{F}^T) + o_p(1),$$

showing that, $n^{1/2}(\widehat{U} - I_p) = O_p(1)$. Consequently, also $n^{1/2}(\widehat{\Gamma} - I_p) = O_p(1)$.

Finally, we next proceed as in the proof of ii) in Theorem 2 of Miettinen et al. (2016) to obtain that, as $n \rightarrow \infty$,

$$n^{1/2}(\widehat{\Gamma}_{i,i} - 1) = -(1/2)n^{1/2}(\widehat{D}_{0,i,i} - 1) + o_p(1)$$

and for $i \neq j$,

$$n^{1/2}\widehat{\Gamma}_{i,j} = \frac{\sum_{l=1}^k \{D(f_l)_i - D(f_l)_j\} n^{1/2} \{\widehat{D}(f_l)_{i,j} - D(f_l)_i \widehat{D}_{0,i,j}\}}{\sum_{r=1}^k \{D(f_r)_i - D(f_r)_j\}^2} + o_p(1).$$

Hence, the lemma follows from the definition of $A_0, A_1, \dots, A_k, B$. $\square$

LEMMA B.7. Assume the same settings and conditions as in Lemma B.5. Let $\Sigma(f, g)$ be as in Proposition B.3 with R replaced by $\text{cov}(z)$ where z is the $np \times 1$ vector defined by, for $i = 1, \dots, n$, $j = 1, \dots, p$, $z_{(i-1)p+j} = Z_j(s_i)$. Let $f_0(x) = I(x=0)$. Let $\tilde{V}(f_1, \dots, f_k)$ be the $(k+1)p^2 \times (k+1)p^2$ matrix, composed of $(k+1)^2$ blocks of sizes $p^2 \times p^2$ with block $(i+1), (j+1)$ equal to $\Sigma(f_i, f_j)$ for $i, j = 0, \dots, p$.

Let G be the $p^2 \times (k+1)p^2$ matrix defined by $G = B(A_0, A_1, \dots, A_k)$, for $n \geq n_0$, with the notation of Lemma B.6. Let $M_{\Omega^{-1}}$ be as in Proposition B.4 and let, for $n \geq n_0$,

$$F = M_{\Omega^{-1}} G \tilde{V}(f_1, \dots, f_k) G^T M_{\Omega^{-1}}^T.$$

Then, $\widehat{\Gamma} = \widehat{\Gamma}\{\widehat{M}_0, \widehat{M}(f_1), \dots, \widehat{M}(f_k)\}$ satisfying (A9), with $\widehat{D}_0, \widehat{D}(f_1), \dots, \widehat{D}(f_k)$ replaced by $\widehat{M}_0, \widehat{M}(f_1), \dots, \widehat{M}(f_k)$, can be chosen so that, with Q_n the distribution of $n^{1/2}\text{vect}(\widehat{\Gamma} - \Omega^{-1})$, we have as $n \rightarrow \infty$

$$d_w\{Q_n, \mathcal{N}(0, F)\} \rightarrow 0.$$

Proof. The proof is the same as that of Proposition B.4. In particular, for $\widehat{\Gamma}\{\widehat{D}_0, \widehat{D}(f_1), \dots, \widehat{D}(f_k)\}$ satisfying (A9), the matrix $\widehat{\Gamma}\{\widehat{D}_0, \widehat{D}(f_1), \dots, \widehat{D}(f_k)\}\Omega^{-1}$ satisfies (A9), with $\widehat{D}_0, \widehat{D}(f_1), \dots, \widehat{D}(f_k)$ replaced by $\widehat{M}_0, \widehat{M}(f_1), \dots, \widehat{M}(f_k)$. $\square$

PROPOSITION B.6. Assume the same settings and conditions as in Lemma B.5. Let $(\widehat{\Gamma}_n)_{n \in \mathbb{N}}$ be any sequence of $p \times p$ matrices so that for any $n \in \mathbb{N}$, $\widehat{\Gamma}_n = \widehat{\Gamma}_n\{\widehat{M}_0, \widehat{M}(f_1), \dots, \widehat{M}(f_k)\}$ satisfies (A9), with $\widehat{D}_0, \widehat{D}(f_1), \dots, \widehat{D}(f_k)$ replaced by $\widehat{M}_0, \widehat{M}(f_1), \dots, \widehat{M}(f_k)$. Then, there exists a sequence of permutation matrices (P_n) and a sequence of diagonal matrices (D_n) , with diagonal components in $\{-1, 1\}$ so that, with $\tilde{\Gamma}_n = D_n P_n \widehat{\Gamma}_n$, the sequence $(\tilde{\Gamma}_n)$ satisfies the conclusions of Lemma B.5, with the limit I_p replaced by Ω^{-1} .

Proof. With the notation of the proof of Lemma B.5, for $\widehat{\Gamma}_n$ satisfying (A9), there exist P_n, D_n , as described in the proposition, so that $D_n P_n \widehat{\Gamma}_n \widehat{D}_0^{1/2} \in E_0$ and $D_n P_n \widehat{\Gamma}_n$ satisfies (A9). Hence, with the same argument as in the proof of the last part of Lemma B.5, we have $D_n P_n \widehat{\Gamma}_n \rightarrow I_p$ in probability as $n \rightarrow \infty$. Furthermore, as in the proof of Lemma B.7, we can show that $D_n P_n \widehat{\Gamma}_n \Omega^{-1}$ satisfies the conclusion of this lemma. The proof is concluded by observing that any matrix $\tilde{\Gamma}$ satisfies (A9), with $\widehat{D}_0, \widehat{D}(f_1), \dots, \widehat{D}(f_k)$ replaced by $\widehat{M}_0, \widehat{M}(f_1), \dots, \widehat{M}(f_k)$, if and only if the corresponding matrix $\tilde{\Gamma}\Omega$ satisfies (A9). $\square$

PROPOSITION B.7. Assume the same settings and conditions as in Lemma B.5. Let $(\widehat{\Gamma}_n)_{n \in \mathbb{N}}$ be any sequence of $p \times p$ matrices so that for any $n \in \mathbb{N}$, $\widehat{\Gamma}_n = \widehat{\Gamma}_n\{\widehat{M}_0, \widehat{M}(f_1), \dots, \widehat{M}(f_k)\}$ satisfies (A9), with $\widehat{D}_0, \widehat{D}(f_1), \dots, \widehat{D}(f_k)$ replaced by $\widehat{M}_0, \widehat{M}(f_1), \dots, \widehat{M}(f_k)$. Then, there exists a sequence of permutation matrices (P_n) and a sequence of diagonal matrices (D_n) , with diagonal components in $\{-1, 1\}$ so that, with $\tilde{\Gamma}_n = D_n P_n \widehat{\Gamma}_n$, the sequence $(\tilde{\Gamma}_n)$ satisfies the conclusions of Lemma B.7.

Proof. The proof is the same as that of Proposition B.6. $\square$

The results of Propositions 6 and 7 derive directly from Propositions B.6 and B.7.

LEMMA B.8. *Let Conditions A1 and A2 hold. Let f satisfy Condition A3. Let $\bar{X} = n^{-1} \sum_{i=1}^n X(s_i)$. Let*

$$\widehat{M}_{\text{st}}(f) = n^{-1} \sum_{i=1}^n \sum_{j=1}^n f(s_i - s_j) \{X(s_i) - \bar{X}\} \{X(s_j) - \bar{X}\}^T.$$

Then as $n \rightarrow \infty$

$$\widehat{M}_{\text{st}}(f) - \widehat{M}(f) = O_p(n^{-1}).$$

Proof. Let $a, l \in \{1, \dots, p\}$ and let $f_{i,j} = f(s_i - s_j)$. We have

$$\begin{aligned} \widehat{M}_{\text{st}}(f)_{a,l} - \widehat{M}(f)_{a,l} &= n^{-1} \sum_{i=1}^n \sum_{j=1}^n f_{i,j} \{X_a(s_i)X_l(s_j) - X_a(s_i)\bar{X}_l - \bar{X}_a X_l(s_j) + \bar{X}_a \bar{X}_l\} \\ &\quad - n^{-1} \sum_{i=1}^n \sum_{j=1}^n f_{i,j} X_a(s_i)X_l(s_j) \\ &= -\bar{X}_l \{n^{-1} \sum_{i=1}^n \sum_{j=1}^n f_{i,j} X_a(s_i)\} - \bar{X}_a \{n^{-1} \sum_{i=1}^n \sum_{j=1}^n f_{i,j} X_l(s_j)\} \\ &\quad + n^{-1} \sum_{i=1}^n \sum_{j=1}^n f_{i,j} \bar{X}_a \bar{X}_l. \end{aligned} \tag{A15}$$

Now, for $q = 1, \dots, p$, $E(\bar{X}_q) = 0$ and

$$\text{var}(\bar{X}_q) = \frac{1}{n^2} \sum_{i=1}^n \sum_{j=1}^n \text{cov}\{X_q(s_i), X_q(s_j)\}.$$

Also, $\max_{i=1, \dots, n} \sum_{j=1}^n |\text{cov}\{X_q(s_i), X_q(s_j)\}|$ is bounded because of (A1) and Lemma 4 in Furrer et al. (2016). Hence $\text{var}(\bar{X}_q) = O(1/n)$ and so $\bar{X}_q = O_p(n^{-1/2})$.

Also, let

$$\epsilon_q = n^{-1} \sum_{i=1}^n \sum_{j=1}^n f_{i,j} X_q(s_i).$$

Then $E(\epsilon_q) = 0$ and

$$\text{var}(\epsilon_q) = n^{-2} \sum_{i=1}^n \sum_{b=1}^n \left(\sum_{j=1}^n f_{i,j} \right) \left(\sum_{j=1}^n f_{b,j} \right) \text{cov}\{X_q(s_i), X_q(s_b)\}.$$

From Condition A3 and Lemma 4 in Furrer et al. (2016), there exists a finite constant H so that

$$\max_{i=1, \dots, n} \sum_{j=1}^n |f_{i,j}| \leq H.$$

Hence

$$\begin{aligned} \text{var}(\epsilon_q) &\leq H^2 n^{-2} \sum_{i=1}^n \sum_{b=1}^n |\text{cov}\{X_q(s_i), X_q(s_b)\}| \\ &= O(n^{-1}) \end{aligned}$$

as before. Hence $\epsilon_q = O_p(n^{-1/2})$. Also, we have seen above that

$$n^{-1} \sum_{i=1}^n \sum_{j=1}^n |f_{i,j}|$$

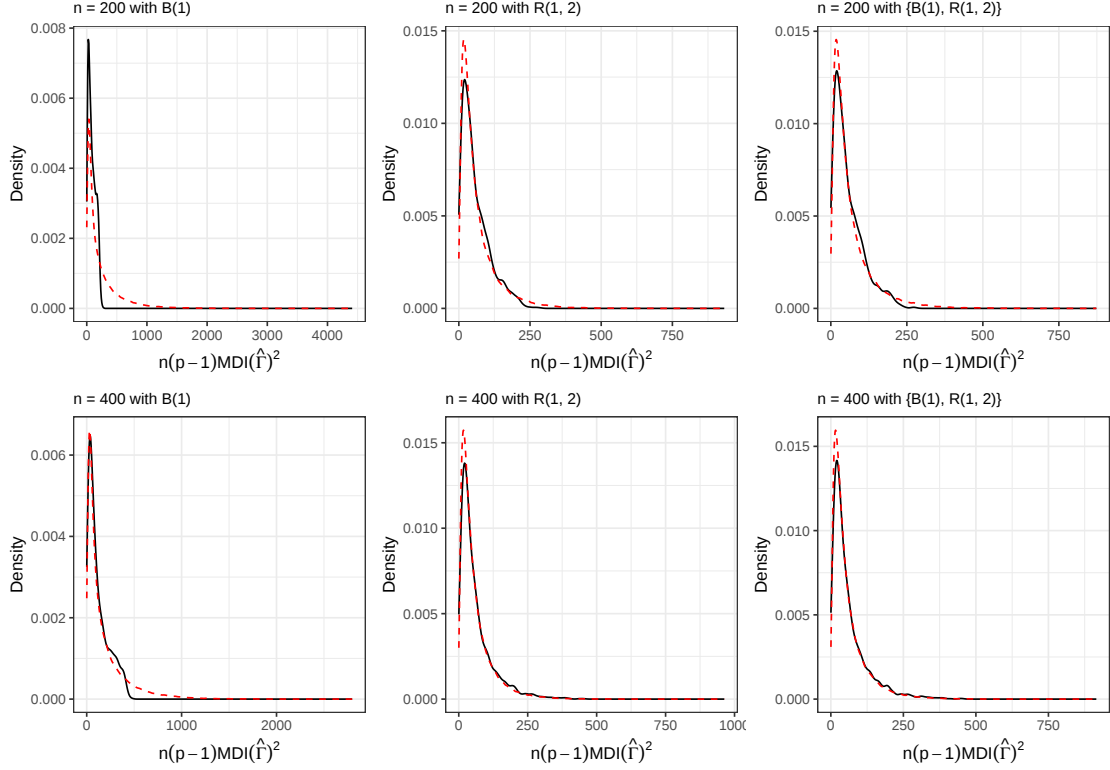


Fig. C.1: The density of $n(p-1)\text{MDI}(\hat{\Gamma})^2$ over 2000 replications, in solid lines, against the density of its asymptotic approximation, in dashed lines, for different combinations of sample size and local covariance matrices.

is bounded. Hence, from (A15), we conclude the proof of the lemma. $\square$

C. SIMULATION COMPLEMENTS

C.1. Asymptotic approximate distribution of the unmixing matrix estimator

In Fig. 3 of Section 5.2, we used only the expected value of the asymptotic approximation. In Fig. C.1 we have plotted the estimated densities of $n(p-1)\text{MDI}(\hat{\Gamma})^2$, in solid lines, against the densities of the corresponding asymptotic approximations, in dashed lines, for all local covariance matrices and a few selected sample sizes. The density functions of the asymptotic approximations are estimated from a sample of 100,000 random variables drawn from the corresponding distributions. Overall, the two densities fit each other rather well, especially for the local covariance matrices involving the ring kernel. This shows that the asymptotic distribution of $n(p-1)\text{MDI}(\hat{\Gamma})^2$ is a good approximation of the true distribution already for small sample sizes.

C.2. A simulation study on individual component estimation accuracy

In this simulation, we investigate the estimation accuracies of the individual latent fields under the spatial blind source separation model.

We use the same setting as in the first simulation study in Section 5. That is, let $X(s) = \Omega Z(s)$ where each of the three independent latent fields has a Matérn covariance function with shape and range parameters $(\kappa, \phi) \in \{(6, 1.2), (1, 1.5), (0.25, 1)\}$, illustrated in the left panel in Fig. 1. The location pattern is taken to be the same growing pattern of nested squares depicted on the left of Fig. 2.

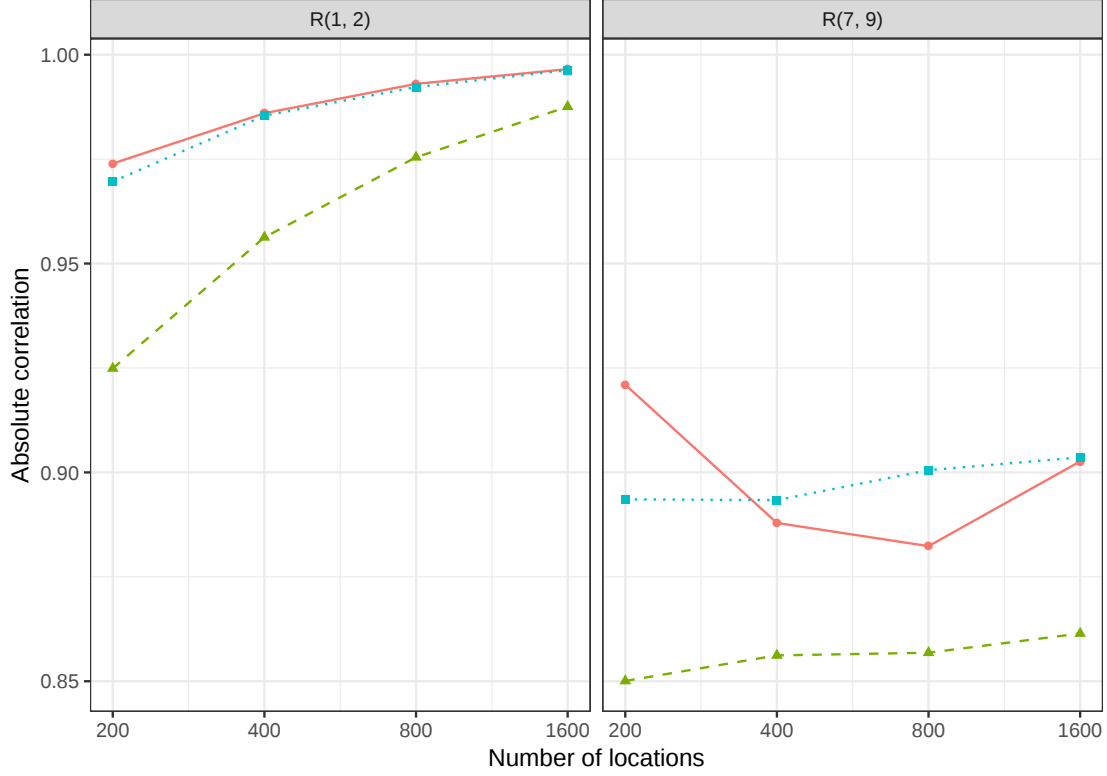


Fig. C.2: Maximal absolute correlations for the first (solid red line), second (dashed green line) and third (dotted blue line) source fields under two kernels and a range of sample sizes.

To quantify the estimation accuracies of the individual components, we use the same strategy as in the real data example of Section 6. Let $\hat{Z}(s) = \{\hat{Z}_1(s), \hat{Z}_2(s), \hat{Z}_3(s)\}^T = \hat{\Gamma}X(s)$ contain the estimated components for a single repetition of the simulation. For each of the three true sources, $j = 1, 2, 3$, we record the maximum absolute sample correlation between $Z_j(s)$ and $\hat{Z}_l(s)$ over $l = 1, 2, 3$. The larger the maximum absolute correlation, the better the source field j was estimated.

Due to the affine equivariance of the estimators, the estimated components $\hat{\Gamma}X(s)$ are invariant to the choice of Ω , up to their signs and order. More precisely, let $\hat{\Gamma}(I_p)$ be computed from $\{Z(s_i)\}_{i=1,\dots,n}$ according to (5), let $\hat{Z}_{I_p}(s) = \hat{\Gamma}(I_p)Z(s)$ and recall that $\hat{\Gamma}$ is computed from $\{X(s_i)\}_{i=1,\dots,n}$ according to (5). Then we have

$$\hat{Z}(s) = \hat{\Gamma}X(s) = \hat{\Gamma}\Omega Z(s) = \hat{\Gamma}(I_p)\Omega^{-1}\Omega Z(s) = \hat{\Gamma}(I_p)Z(s) = \hat{Z}_{I_p}(s),$$

up to order and signs of the components. Therefore, it is without loss of generality that we may again assume that $\Omega = I_3$. Any other choice of Ω would lead to exactly the same results.

The mean maximum source-wise absolute correlations over 2000 repetitions are shown in Fig. C.2 for a range of sample sizes. We have used two choices of kernels, $R(1, 2)$ and $R(7, 9)$. The first one was chosen due to its good performance in the main simulation study and the second one to see how the individual components are estimated under a bad choice of kernel.

The results indicate that the first and third sources are estimated almost equally well, but that the second source is somewhat more difficult to estimate. We postulate that this is caused by its corresponding covariance function being the middle one in the left panel in Fig. 1. That is, the first and third sources are unlikely to be mixed with each other due to their extremal covariance functions. The second source, on the other hand, is between the other two and can be mistaken for both the first and the third source. We

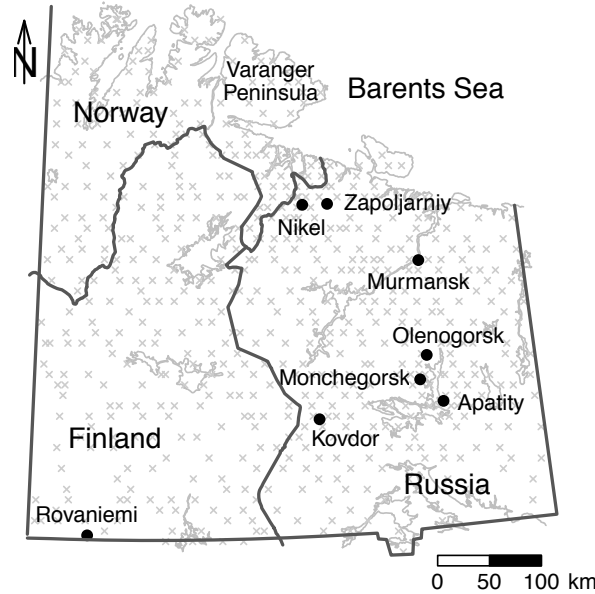


Fig. D.1: Sampling area of the Kola data. The gray crosses indicate sampling locations.

also note in the right panel of Fig. C.2 that using an inferior choice of a kernel leads to an overall worse estimation accuracy for all sources. Finally, on the left-hand side of Fig. C.2, we observe a convergence to one of the maximum absolute correlation, under the appropriate choice of kernel $R(1, 2)$.

D. FURTHER DETAILS CONCERNING THE REAL DATA EXAMPLE

Figure D.1 describes the sampling area from the Kola data and Figures D.2-D.4 visualize the components discussed in Section 6.

REFERENCES

- BACHOC, F. (2014). Asymptotic analysis of the role of spatial sampling for covariance parameter estimation of Gaussian processes. *J. Multivariate Anal.* **125**, 1–35.
- BELOUCHRANI, A., ABED-MERAIM, K., CARDOSO, J.-F. & MOULINES, E. (1997). A blind source separation technique using second-order statistics. *IEEE T. Signal Proces.* **45**, 434–444.
- CLARKSON, D. B. (1988). Remark AS R71: A remark on algorithm as 211. the F-G diagonalization algorithm. *J. R. Stat. Soc. C-Appl.* **37**, 147–151.
- COMON, P. & JUTTEN, C. (2010). *Handbook of Blind Source Separation: Independent component analysis and applications*. Academic press.
- CRESSIE, N. (1993). *Statistics for Spatial Data*. John Wiley & Sons, Inc., New York, 2nd ed.
- DE IACO, S., MYERS, D. E., PALMA, M. & POSA, D. (2013). Using simultaneous diagonalization to identify a space-time linear coregionalization model. *Math. Geosci.* **45**, 69–86.
- DUDLEY, R. M. (2002). *Real Analysis and Probability*. Cambridge University Press.
- EDDELBUETTEL, D. & SANDERSON, C. (2014). Repparmadillo: Accelerating R with high-performance C++ linear algebra. *Comput. Stat. Data. An.* **71**, 1054–1063.
- EMERY, X. (2010). Iterative algorithms for fitting a linear model of coregionalization. *Comput. Geosci.* **36**, 1150–1160.
- FILZMOSER, P. (2015). *StatDA: Statistical Analysis for Environmental Data*. R package version 1.6.9.
- FURRER, R., BACHOC, F. & DU, J. (2016). Asymptotic properties of multivariate tapering for estimation and prediction. *J. Multivariate Anal.* **149**, 177–191.

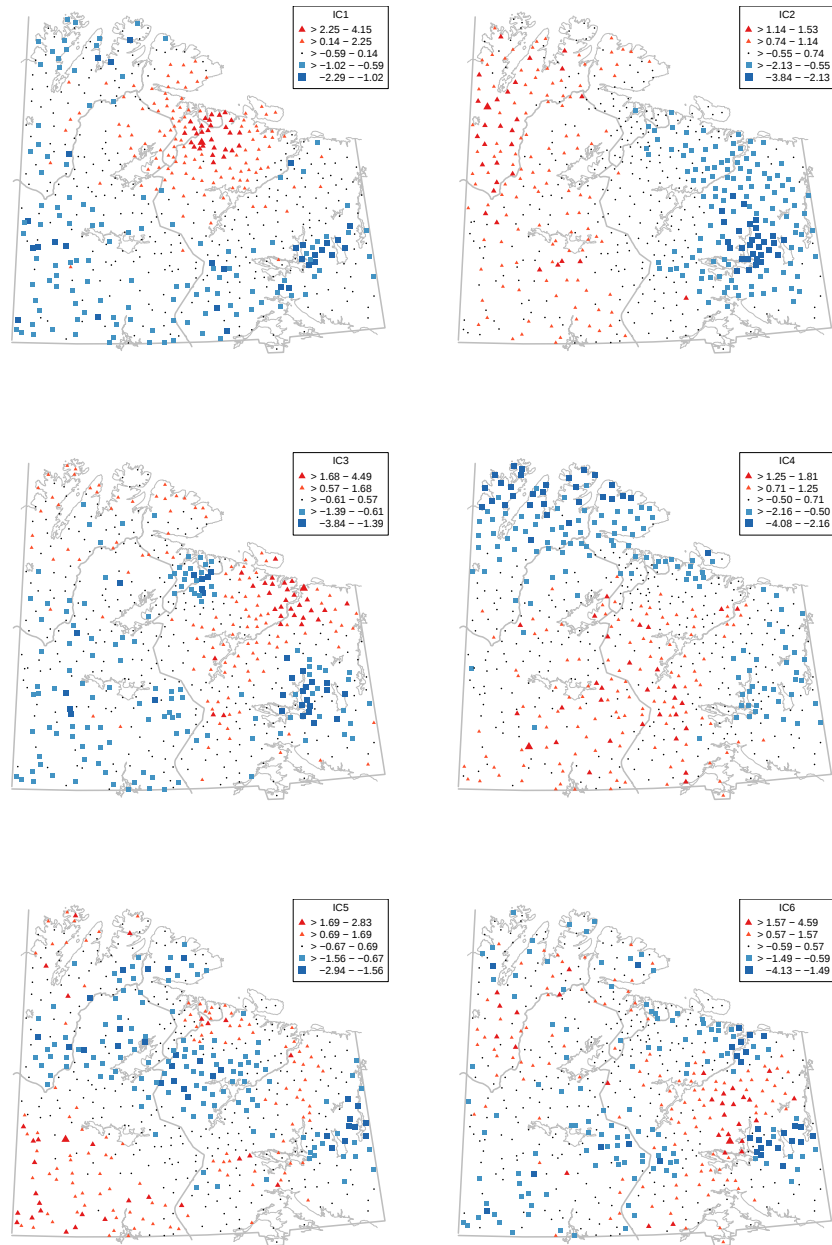


Fig. D.2: The six most interesting components using the covariance matrix and $B(50)$. Used as gold standard following Nordhausen et al. (2015).

- GELFAND, A. E., SCHMIDT, A. M., BANERJEE, S. & SIRMANS, C. F. (2004). Nonstationary multivariate process modeling through spatially varying coregionalization. *TEST* **13**, 263–312.
- GENTON, M. G. & KLEIBER, W. (2015). Cross-covariance functions for multivariate geostatistics. *Stat. Sci.* **30**, 147–163.
- GOULARD, M. & VOLTZ, M. (1992). Linear coregionalization model: Tools for estimation and choice of cross-variogram matrix. *Math. Geol.* **24**, 269–286.

- ILLNER, K., MIETTINEN, J., FUCHS, C., TASKINEN, S., NORDHAUSEN, K., OJA, H. & THEIS, F. J. (2015). Model selection using limiting distributions of second-order blind source separation algorithms. *Signal Process.* **113**, 95–103.
- ILMONEN, P., NEVALAINEN, J. & OJA, H. (2010a). Characteristics of multivariate distributions and the invariant coordinate system. *Stat. Probabil. Lett.* **80**, 1844–1853.
- ILMONEN, P., NORDHAUSEN, K., OJA, H. & OLLILA, E. (2010b). A new performance index for ICA: properties, computation and asymptotic analysis. In *Latent Variable Analysis and Signal Separation*. Springer, pp. 229–236.
- ISTAS, J. & LANG, G. (1997). Quadratic variations and estimation of the local Hölder index of a gaussian process. *Ann. I. H. Poincaré-Pr.* **33**, 407–436.
- MATHAI, A. M. & PROVOST, S. B. (1992). *Quadratic forms in random variables: theory and applications*. Dekker.
- MATILAINEN, M., NORDHAUSEN, K. & OJA, H. (2015). New independent component analysis tools for time series. *Stat. Probabil. Lett.* **105**, 80–87.
- MIETTINEN, J., ILLNER, K., NORDHAUSEN, K., OJA, H., TASKINEN, S. & THEIS, F. (2016). Separation of uncorrelated stationary time series using autocovariance matrices. *J. Time Ser. Anal.* **37**, 337–354.
- MIETTINEN, J., NORDHAUSEN, K., OJA, H. & TASKINEN, S. (2012). Statistical properties of a blind source separation estimator for stationary time series. *Stat. Probabil. Lett.* **82**, 1865–1873.
- MIETTINEN, J., NORDHAUSEN, K., OJA, H. & TASKINEN, S. (2014). Deflation-based separation of uncorrelated stationary time series. *J. Multivariate Anal.* **123**, 214–227.
- MIETTINEN, J., NORDHAUSEN, K. & TASKINEN, S. (2017). Blind source separation based on joint diagonalization in R: The packages JADE and BSSasyp. *J. Stat. Softw.* **76**, 1–31.
- NORDHAUSEN, K. (2014). On robustifying some second order blind source separation methods for nonstationary time series. *Stat. Pap.* **55**, 141–156.
- NORDHAUSEN, K. & OJA, H. (2018). Independent component analysis: A statistical perspective. *WIREs Comput. Stat.* **10**, e1440.
- NORDHAUSEN, K., OJA, H., FILZMOSER, P. & REIMANN, C. (2015). Blind source separation for spatially correlated compositional data. *Math. Geosci.* **47**, 753–770.
- R CORE TEAM (2019). *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing, Vienna, Austria.
- REIMANN, C., FILZMOSER, P., GARRETT, R. & DUTTER, R. (2008). *Statistical Data Analysis Explained. Applied Environmental Statistics with R*. Chichester: Wiley.
- RIBEIRO JR, P. J. & DIGGLE, P. J. (2016). *geoR: Analysis of Geostatistical Data*. R package version 1.7-5.2.
- SUN, D. & SUN, J. (2002). Strong semismoothness of eigenvalues of symmetric matrices and its application to inverse eigenvalue problems. *SIAM J. Numer. Anal.* **40**, 2352–2367.
- TANG, A. C., LIU, J.-Y. & SUTHERLAND, M. T. (2005). Recovery of correlated neuronal sources from EEG: the good and bad ways of using SOBI. *Neuroimage* **28**, 507–519.
- TASKINEN, S., MIETTINEN, J. & NORDHAUSEN, K. (2016). A more efficient second order blind identification method for separation of uncorrelated stationary time series. *Stat. Probabil. Lett.* **116**, 21–26.
- VIRTA, J., LIETZÉN, N., ILMONEN, P. & NORDHAUSEN, K. (2018). Asymptotically and computationally efficient tensorial JADE. *arXiv preprint arXiv:1808.00791*.
- VIRTA, J. & NORDHAUSEN, K. (2019). Determining the signal dimension in second order source separation. *To appear in Stat. Sinica*.
- WACKERNAGEL, H. (2003). *Multivariate Geostatistics: An Introduction with Applications*. Berlin: Springer-Verlag, 3rd ed.
- XIE, T. & MYERS, D. E. (1995). Fitting matrix-valued variogram models by simultaneous diagonalization (Part I: Theory). *Math. Geol.* **27**, 867–875.
- XIE, T., MYERS, D. E. & LONG, A. E. (1995). Fitting matrix-valued variogram models by simultaneous diagonalization (Part II: Application). *Math. Geol.* **27**, 877–888.
- ZHANG, H. (2007). Maximum-likelihood estimation for multivariate spatial linear coregionalization models. *Environmetrics* **18**, 125–139.

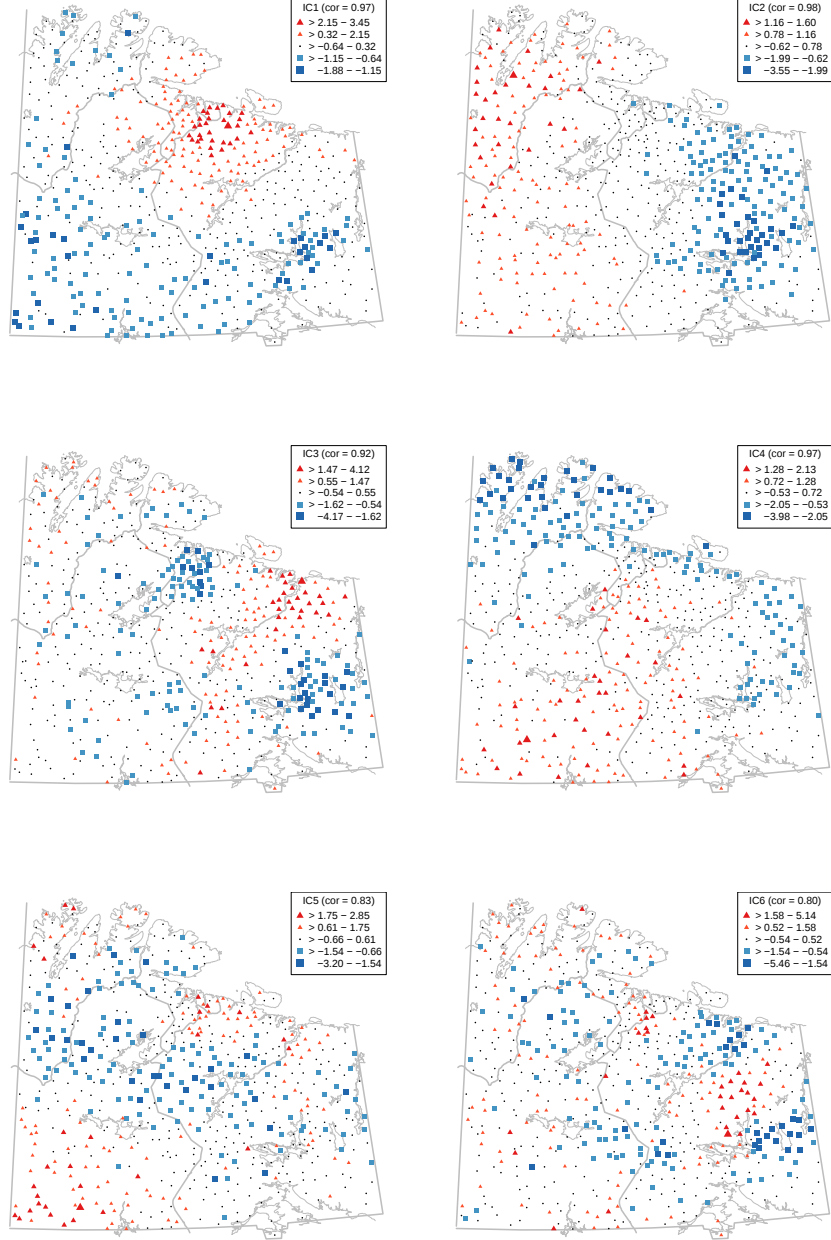


Fig. D.3: The six components with highest absolute correlations with the components from Figure D.2 when jointly diagonalising the covariance matrix and the ring kernels $R(0, 25)$, $R(25, 50)$, $R(50, 75)$, $R(75, 100)$. The correlations to the corresponding components of Fig. D.2 are given in the legends.

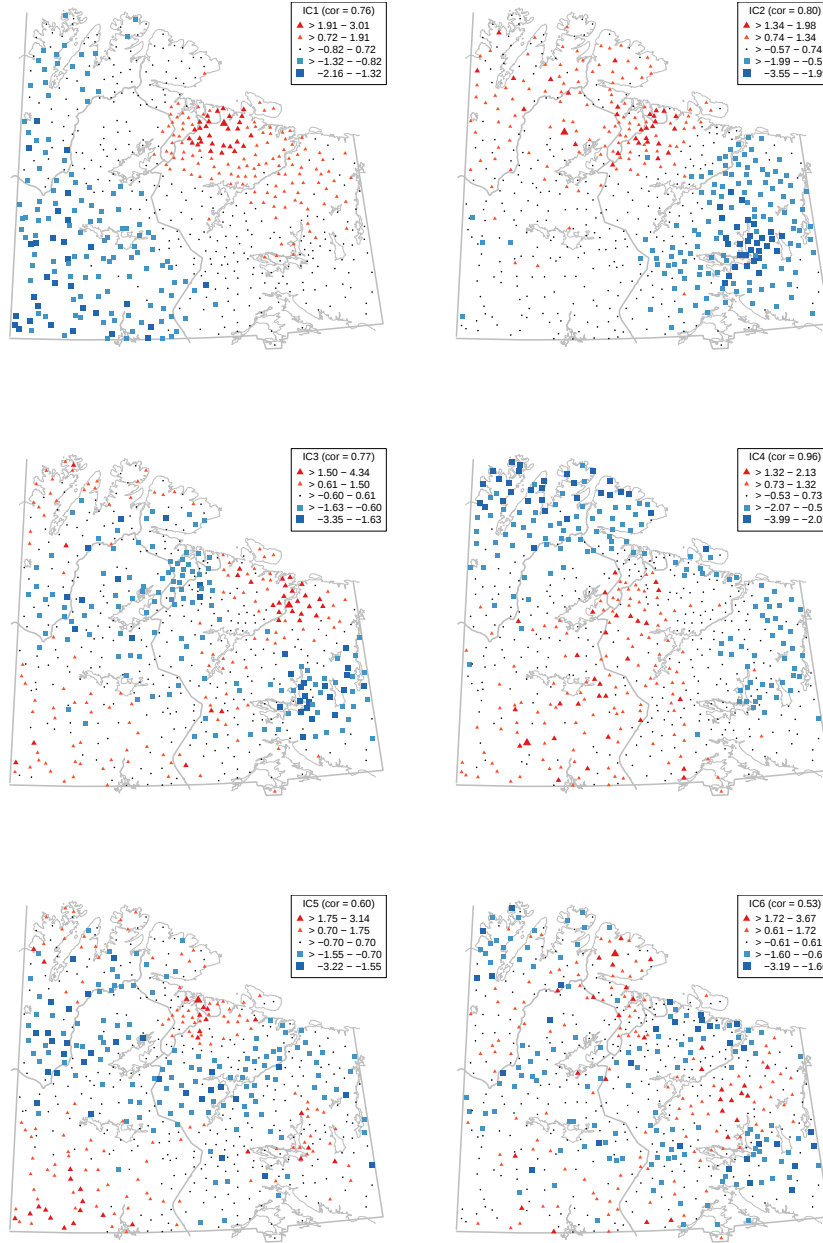


Fig. D.4: The six components with highest absolute correlations with the components from Fig. D.2 when jointly diagonalising the covariance matrix and $B(100)$. The correlations to the corresponding components of Fig. D.2 are given in the legends.