



Automated data-driven selection of the hyperparameters for Total-Variation based texture segmentation

Barbara Pascal, Samuel Vaiter, Nelly Pustelnik, Patrice Abry

► To cite this version:

Barbara Pascal, Samuel Vaiter, Nelly Pustelnik, Patrice Abry. Automated data-driven selection of the hyperparameters for Total-Variation based texture segmentation. *Journal of Mathematical Imaging and Vision*, 2021, 63, pp.923-952. 10.1007/s10851-021-01035-1 . hal-03044181

HAL Id: hal-03044181

<https://hal.science/hal-03044181>

Submitted on 7 Dec 2020

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Automated data-driven selection of the hyperparameters for Total-Variation based texture segmentation.*

Barbara Pascal[†] Samuel Vaiter[‡] Nelly Pustelnik[†]
 Patrice Abry[†]

April 2020

1 Introduction

Numerous problems in signal and image processing consist in finding the best possible estimate $\hat{\mathbf{x}}$ of a quantity $\mathbf{x} \in \mathcal{H}$ from an observation $\mathbf{y} \in \mathcal{G}$ (where $\mathcal{H}$ and $\mathcal{G}$ are Hilbert spaces isomorphic to $\mathbb{R}^N$ and $\mathbb{R}^P$ respectively), potentially corrupted by a linear operator $\Phi : \mathcal{H} \rightarrow \mathcal{G}$, which encapsulates deformation or information loss, and by some additive zero-mean Gaussian noise $\zeta \sim \mathcal{N}(\mathbf{0}_P, \mathcal{S})$, with *known* covariance matrix $\mathcal{S} \in \mathbb{R}^{P \times P}$, leading to the general observation model

$$\mathbf{y} = \Phi \bar{\mathbf{x}} + \zeta. \quad (1)$$

Examples resorting to inverse problems include image restoration [13, 55], inpainting [17], texture-geometry decomposition [3], but also texture segmentation as recently proposed in [51]. A widely investigated path for the estimation of underlying $\bar{\mathbf{x}}$ is linear regression [10, 41], providing an unbiased linear regression estimator $\hat{\mathbf{x}}_{\text{LR}}$. Yet corresponding estimates suffer from large variances, which can lead to dramatic errors in the presence of noise ζ [6].

An alternative relies on the construction of *parametric estimators*

$$\begin{aligned} \mathcal{G} \times \mathbb{R}^L &\longrightarrow \mathcal{H} \\ (\mathbf{y}, \Lambda) &\longmapsto \hat{\mathbf{x}}(\mathbf{y}; \Lambda) \end{aligned} \quad (2)$$

allowing some estimation bias, and thus leading to drastic decrease of the variance. Given some prior knowledge about ground truth $\bar{\mathbf{x}}$, e.g. [35], either sparsity of the variable $\bar{\mathbf{x}}$ [63], of its derivative [37, 59, 66] or of its wavelet transform [27], one can build *parametric estimators* performing a compromise

*Work supported by Defi Imag'in SIROCCO and by ANR-16-CE33-0020 MultiFracs, France and by ANR GraVa ANR-18-CE40-0005.

[†]Univ Lyon, ENS de Lyon, Univ Lyon 1, CNRS, Laboratoire de Physique, F-69342 Lyon, France (firstname.lastname@ens-lyon.fr).

[‡]CNRS & Université de Bourgogne Franche-Comté, Dijon, France. (samuel.vaiter@u-bourgogne.fr)

between fidelity to the model (1) and structure constraints on the estimation. In general, the compromise is tuned by a small number $L = \mathcal{O}(1)$ parameters, stored in a vector $\mathbf{\Lambda} \in \mathbb{R}^L$. A very popular class of *parametric estimators* relies on a penalization of a least squares data fidelity term formulated as a minimization problem

$$\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}) \in \underset{\mathbf{x} \in \mathcal{H}}{\text{Argmin}} \|\mathbf{y} - \Phi \mathbf{x}\|_{\mathcal{W}}^2 + \|\mathbf{U}_{\mathbf{\Lambda}} \mathbf{x}\|_q^q \quad (3)$$

with $\|\cdot\|_{\mathcal{W}}$ the Mahalanobis distance associated with $\mathcal{W} \in \mathbb{R}^{P \times P}$, defined as

$$\|\mathbf{y} - \Phi \mathbf{x}\|_{\mathcal{W}} \triangleq \sqrt{(\mathbf{y} - \Phi \mathbf{x})^\top \mathcal{W} (\mathbf{y} - \Phi \mathbf{x})}. \quad (4)$$

$\mathbf{U}_{\mathbf{\Lambda}} : \mathcal{H} \rightarrow \mathcal{Q}$ is a linear operator parametrized by $\mathbf{\Lambda}$ and $\|\cdot\|_q$ the ℓ_q -norm with $q \geq 1$ in Hilbert space $\mathcal{Q}$.

Least Squares. While Ordinary Least Squares involve usual ℓ_2 squared norm as data-fidelity term, that is $\mathcal{W} = \mathbf{I}_P$, Generalized Least Squares [61] make use of the covariance structure of the noise through $\mathcal{W} = \mathcal{S}^{-1}$, encapsulating all the observation statistics in the case of Gaussian noise. This generalized approach is equivalent to decorrelating the data and equalizing noise levels before performing the regression. Further, the Gauss-Markov theorem [1] asserts that minimizing Weighted Least Squares provides the best linear estimator of $\mathbf{x}$, advocating for the use of Mahalanobis distance as data fidelity term in penalized Least Squares. Yet, in practice, Generalized Least Squares (or Weighted Least Squares in the case when $\mathcal{S}$ is diagonal) requires not only the knowledge of the covariance matrix, but also to be able to invert it. For uncorrelated data, $\mathcal{S}$ is diagonal and, provided that it is well-conditioned, it is easy to invert numerically. On the contrary, computing $\mathcal{S}^{-1}$ might be extremely challenging for correlated data since $\mathcal{S}$ is not diagonal anymore and has a size scaling like the square of the dimension of $\mathcal{G}$. Thus, to handle possibly correlated Gaussian noise ζ , using Ordinary Least Squares is often mandatory, even though it does not benefit from same theoretical guarantees that Generalized Least Squares. Nevertheless, we will show that the knowledge of $\mathcal{S}$ is far from being useless, since it is possible to take advantage of it when estimating the quadratic risk.

Penalization. Appropriate choice of q and $\mathbf{U}_{\mathbf{\Lambda}}$ covers a large variety of well-known estimators. Linear filtering is obtained for $q = 2$ [31], the shape of the filter being encapsulated in operator $\mathbf{U}_{\mathbf{\Lambda}}$ [37], the hyperparameters $\mathbf{\Lambda}$ tuning e.g. its band-width. It is very common in image processing to impose priors on the spatial gradients of the image, using the finite discrete horizontal and vertical difference operator $\mathbf{D}$ and one regularization parameter $\mathbf{\Lambda} = \lambda > 0$ ($L = 1$). For example, smoothness of the estimate is favored using ℓ_2 squared norm, performing Tikhonov regularization [37, 66], in which $q = 2$ and $\|\mathbf{U}_{\mathbf{\Lambda}} \mathbf{x}\|_q^q \triangleq \lambda \|\mathbf{D} \mathbf{x}\|_2^2$. Another standard penalization is the anisotropic total variation [59] $\|\mathbf{U}_{\mathbf{\Lambda}} \mathbf{x}\|_q^q \triangleq \lambda \|\mathbf{D} \mathbf{x}\|_1$, corresponding to $q = 1$, where the ℓ_1 -norm enforces sparsity of spatial gradients.

Risk estimation. The purpose of Problem (3) is to obtain a faithful estimation $\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ of ground truth $\mathbf{x}$, the error being measured by the so-called *quadratic*

risk

$$\mathbb{E} \|\mathbf{B}\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}) - \mathbf{B}\bar{\mathbf{x}}\|_{\mathcal{W}}^2 \quad (5)$$

with $\mathbf{B}$ a linear operator, which enables to consider various types of risk. For instance, when $\mathbf{B} = \mathbf{\Pi}$ is a projector on a subset of $\mathcal{H}$ [30], the *projected quadratic risk* (5) measures the estimation error on the projected quantity $\mathbf{\Pi}\bar{\mathbf{x}}$. This case includes the usual quadratic risk when $\mathbf{B} = \mathbf{I}_N$. Conversely, when $\mathbf{B} = \mathbf{\Phi}$, the risk (5) quantifies the quality of the prediction $\hat{\mathbf{y}}(\mathbf{y}; \mathbf{\Lambda}) \triangleq \mathbf{\Phi}\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ with respect to the noise-free observation $\bar{\mathbf{y}} \triangleq \mathbf{\Phi}\bar{\mathbf{x}}$ lying in $\mathcal{G}$, and is known as the *prediction risk*.

The main issue is that one does not have access to ground truth $\bar{\mathbf{x}}$. Hence, measuring the quadratic risk (5) first requires to derive an estimator of

$$\mathbb{E} \|\mathbf{B}\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}) - \mathbf{B}\bar{\mathbf{x}}\|_{\mathcal{W}}^2$$

not involving $\bar{\mathbf{x}}$.

This problem was handled originally in the case of independent, identically distributed, (i.i.d.) Gaussian linear model, that is for scalar covariance matrix $\mathbf{S} = \rho^2 \mathbf{I}_P$, by Stein [60, 64], performing a clever integration by part, leading to Stein's Unbiased Risk Estimate (SURE) [26, 47, 57, 65], initially formulated for the *prediction risk*,

$$\|(\mathbf{\Phi}\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}) - \mathbf{y})\|_{\mathcal{W}}^2 + 2\rho^2 \text{Tr}(\partial_{\mathbf{y}}\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})) - P\rho^2, \quad (6)$$

whose expected value equals quadratic risk (5) with $\mathbf{B} = \mathbf{\Phi}$. In the past years SURE was intensively used both in statistical, signal and image processing applications [11, 26, 53]. It was recently extended to the case of independent but not identically distributed noise [19, 73], corresponding to diagonal covariance matrix $\mathbf{S} = \text{diag}(\sigma_1^2, \dots, \sigma_P^2)$, and to the case when the noise is Gaussian with potential correlations, with very general covariance matrix $\mathbf{S}$. Yet, to the best of our knowledge, very few numerical assessments are available for Gaussian noise with non-scalar covariance matrices. A notable exception is [19], in which numerical experiments are run on uncorrelated multi-component data, the components experiencing different noise levels. The noise being assumed independent, this corresponds to a diagonal covariance matrix $\mathbf{S} = \text{diag}(\rho_1^2, \dots, \rho_P^2)$, with ρ_i^2 the variance of the noise of the i^{th} component.

Further, in the case when the noise is neither independent identically distributed nor Gaussian, Generalized Stein Unbiased Risk Estimators were proposed, e.g. for Exponential Families [30, 38] or Poisson noise [39, 42, 45].

As for practical evaluation of Stein estimator, more sophisticated tools might be required to evaluate the second term of (6), notably when $\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ is obtained from a proximal splitting algorithm [4, 20, 21, 50] solving Problem (3). Indeed Stein estimator involves the Jacobian of $\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ with respect to observations $\mathbf{y}$, which might not be directly accessible in this case. In order to manage this issue, Vonesch *et al* proposed in [69] to perform recursive forward differentiation inside the splitting scheme solving (3), which benefits from few theoretical results from [32]. This approach, even if remaining partially heuristic, proved to be efficient for a large class of problems [24].

Hyperparameter tuning. Equation (3) clearly shows that the estimate $\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ drastically depends on the choice of regularization parameters $\mathbf{\Lambda}$. Thus, fine-tuning of regularization parameters is a long-standing problem in signal and

image processing. A common formulation of this problem consists in minimizing the quadratic risk with respect to regularization parameters $\mathbf{\Lambda}$, solving:

$$\underset{\mathbf{\Lambda}}{\text{minimize}} \mathbb{E} \|\mathbf{B}\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}) - \mathbf{B}\bar{\mathbf{x}}\|_{\mathbf{W}}^2. \quad (7)$$

As emphasized in [30], approximate solution of (7) found selecting among the estimates $(\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}))_{\mathbf{\Lambda} \in \mathbb{R}^L}$ the one reaching lowest SURE (6), as proposed in pioneering work [60], leads to lower mean square error than classical Maximum Likelihood approaches applied to Model (1).

The most direct method solving (7) consists in computing SURE (6) over a grid of parameters [27, 30, 57], and to select the parameter of the grid for which SURE is minimal. Yet, grid search methods suffers from a high computation cost for several reasons. First of all, the size of the grid scaling algebraically with the number of regularization parameters L , exhaustive grid search is often inaccessible. Recently, random strategies were proposed to improve grid search efficiency [7]. Yet, for $L \geq 3$, it remains very challenging if not unfeasible. Further, an additional difficulty might appear in the case when $\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ is obtained from a splitting algorithm solving Problem (3). Indeed when the regularization term $\|\mathbf{U}_{\mathbf{\Lambda}}\mathbf{x}\|_q^q$ is nonsmooth, the proximal algorithms solving (3) suffers from slow convergence rate, making the evaluation of Stein estimator at each point of the grid very time consuming. Although accelerated schemes were proposed [5, 16], grid search with $L \geq 2$ remains very costly, preventing from practical use.

When a closed-form expression of Stein estimator is available, exact function minimization over the regularization parameters $\mathbf{\Lambda}$ might be possible. This is the case for instance for the Tikhonov penalization for which Thompson *et al.* [62], Galatsanos *et al.* in [33], and Desbat *et al.* in [25] took advantage of the linear closed-form expression of $\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ to find the “best” regularization parameter, i.e. to solve (7). Another well-known closed-form expression holds for soft-thresholding, which is widely used for wavelet-shrinkage denoising e.g. [27, 44]. Note that Generalized Cross Validation [36] also makes use of closed-form expression for parameters tuning, but in a slightly different way, working on prediction risk, solving (7) for $\mathbf{B} = \mathbf{\Phi}$. Generalized Cross Validation and Stein-based estimators were compared independently by Li [43], Thompson [62], and Desbat *et al.* in [25]. Further, Bayesian methods were proposed to deal with very large number of hyperparameters $L \gg 1$, among which Sequential Model-Based Optimization (SMBO), providing smart sampling of the hyperparameter domain [8]. Such methods are particularly adapted to machine learning, as they manage huge amount of hyperparameters without requiring knowledge of the gradient of the cost function [9].

In order to go further than (random) sampling methods, elaborated approaches relying on minimization schemes were proposed, requiring sufficiently smooth risk estimator, as well as access to its derivative with respect to $\mathbf{\Lambda}$. From a C^∞ closed-form expression of Poisson Unbiased Risk Estimate, Deledalle *et al.* [23] proposed a Newton algorithm solving (7). Nevertheless, it does not generalize, since it is very rare that one has access to all the derivatives of the risk estimator. In the case when the noise is Gaussian i.i.d., Chaux *et al.* [19] proposed and assessed numerically an empirical descent algorithm for automatic choice of regularization parameter, but with no convergence guarantee. For i.i.d. Gaussian noise and estimators built as the solution of (3), Deledalle *et al.* [24] proposed

sufficient conditions so that $\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ is differentiable with respect to $\mathbf{\Lambda}$, and then derived the differentiability of Stein’s Unbiased Risk Estimate. Further, they elaborated a Stein Unbiased GrAdient estimator of the Risk (SUGAR) with the aim of performing a quasi-Newton descent solving (7) using BFGS strategy. SUGAR proved its efficiency in the automated hyperparameter selection in a spatial-spectral deconvolution method for large multispectral data corrupted by i.i.d. Gaussian noise [2]

Contributions and outline. We propose a Generalized Stein Unbiased GrAdient estimator of the Risk, for the case of Gaussian noise ζ with *any* covariance matrix $\mathbf{S}$, using the framework of Ordinary Least Squares, that is (5) with $\mathbf{W} = \mathbf{I}_P$, enabling to manage different noise levels and correlations in the observed data.

Section 2 revisits Stein’s Unbiased Estimator of the Risk in the particular case of correlated Gaussian noise with covariance matrix $\mathbf{S}$ and derives the Finite Difference Monte Carlo SURE for this framework, extending [24]. Further, we include a projection operator $\mathbf{B} = \mathbf{\Pi}$ making the model versatile enough to fit various applications.

In this context, Finite Difference Monte Carlo SURE is differentiated with respect to regularization parameters leading to Finite Difference Monte Carlo Generalized Stein Unbiased GrAdient estimator of the Risk, whose asymptotic unbiasedness is demonstrated in Section 3. Generalized Stein Unbiased Risk Estimate and Generalized Stein Unbiased GrAdient estimate of the Risk are embedded in a quasi-Newton optimization scheme for automatic parameters tuning, presented in Section 3.3. Moreover, the case of sequential estimators is discussed in Section 3.2.

Then, in Section 4, the entire proposed procedure is particularized to an original application to texture segmentation based on a wavelet (multiscale) estimation of fractal attributes, proposed in [51, 52]. The texture model is cast into the general formulation (1), $\mathbf{y}$ corresponding to a nonlinear multiscale transform of the image to be segmented. Hence the noise ζ presents both inter-scale and intra-scale correlations, leading to a non-diagonal covariance matrix $\mathbf{S}$. Both Stein Unbiased Risk Estimate and Stein Unbiased GrAdient estimate of the Risk are evaluated with a Finite Difference Monte Carlo strategy, all steps of which are made explicit for the texture segmentation problem.

Finally, Section 5 is devoted to exhaustive numerical simulations assessing the performance of the proposed texture segmentation with automatic regularization parameters tuning. We notably emphasize the importance of taking into account the *full* covariance structure into account in Stein-based approaches.

2 Stein Unbiased Risk Estimate (SURE) with correlated noise

This Section details the extension of Stein Unbiased Risk Estimator (6) when $\mathbf{W} = \mathbf{I}_P$ to the case when observations evidence correlated noise, leading to the Finite Difference Monte Carlo Generalized Stein Unbiased Risk Estimator, $\hat{R}_{\nu, \epsilon}(\mathbf{y}; \mathbf{\Lambda} | \mathbf{S})$, defined in (17).

Notations. For a linear operator $\Phi : \mathcal{H} \rightarrow \mathcal{G}$, the *adjoint* operator is denoted Φ^* and characterized by: for every $x \in \mathcal{H}$, and $y \in \mathcal{G}$, $\langle y, \Phi x \rangle = \langle \Phi^* y, x \rangle$. The Jacobian with respect to observations y of a differentiable estimator $\hat{x}(y; \Lambda)$ is denoted $\partial_\Lambda \hat{x}(y; \Lambda)$.

2.1 Observation model

In this work, we consider observations y , supposed to follow Model (1), as stated in Assumption 1 with a degradation operator Φ assumed to be full-rank, as stated in Assumption 2.

Assumption 1 (Gaussianity). The additive noise $\zeta \in \mathcal{G}$ is Gaussian: $\zeta \sim \mathcal{N}(\mathbf{0}_P, \mathcal{S})$, where $\mathbf{0}_P$ is the null vector of $\mathcal{G}$ and $\mathcal{S} \in \mathbb{R}^{P \times P}$ is the covariance matrix of the noise, where $P = \dim(\mathcal{G})$. Thus, the density probability law associated with the model (1) writes

$$y \sim \frac{1}{\sqrt{(2\pi)^P |\det \mathcal{S}|}} \exp \left(-\frac{\|y - \Phi \bar{x}\|_{\mathcal{S}^{-1}}^2}{2} \right). \quad (8)$$

Assumption 2 (Full-rank). The linear operator $\Phi : \mathcal{H} \rightarrow \mathcal{G}$ is full rank, or equivalently $\Phi^* \Phi$ is invertible.

2.2 Estimation problem

Let $\hat{x}(y; \Lambda)$ be a parametric estimator of ground truth $\bar{x} \in \mathcal{H}$, defined in a unique manner from observations $y \in \mathcal{G}$ and hyperparameters $\Lambda \in \mathbb{R}^L$.

Remark 1. For instance, $\hat{x}(y; \Lambda)$ can be the Penalized Ordinary Least Squares estimator, defined in (3). In this case full-rank Assumption 2 ensures the unicity of the minimizer. Nevertheless, we emphasize that Sections 2 and 3 address Problem (7) in a more general framework.

The possibility that the quantity of interest might be a projection of $\bar{x}$ on a the subspace $\mathcal{I}$ of $\mathcal{H}$ is considered. One can think for instance of physics problems, in which only part of variables have a physical interpretation.

Definition 1. The linear operator $\Pi : \mathcal{H} \rightarrow \mathcal{H}$ performs the orthogonal projection on subspace $\mathcal{I}$ capturing relevant information about $\bar{x}$. Moreover, from both Assumption 2 and the projection operator Π , we define the linear operator $\mathbf{A} : \mathcal{G} \rightarrow \mathcal{H}$ as the composition

$$\mathbf{A} \triangleq \Pi (\Phi^* \Phi)^{-1} \Phi^*. \quad (9)$$

The risk is defined as the *projected estimation* error made on the quantity of interest $\Pi \bar{x}$ by the estimator, measured via an ordinary squared ℓ_2 -norm.

$$R[\hat{x}](\Lambda) \triangleq \mathbb{E}_\zeta \|\Pi \hat{x}(y; \Lambda) - \Pi \bar{x}\|_2^2. \quad (10)$$

Remark 2. Another usual definition of the risk involves the inverse of the covariance matrix [30] through a Mahalanobis distance writing

$$R_M[\hat{x}](\Lambda) \triangleq \mathbb{E}_\zeta \|\Pi \hat{x}(y; \Lambda) - \Pi \bar{x}\|_{\mathcal{S}^{-1}}^2. \quad (11)$$

requiring the knowledge of $\mathcal{S}^{-1}$, which might be non-trivial or even inaccessible for correlated noise presenting non-diagonal covariance matrix. Hence our approach uses exclusively ordinary quadratic risk defined in (10). Nevertheless, these two approaches, even though being different, shares interesting common points which will be mentioned briefly in the following (see Remark 3).

The aim of this work is automatic fine-tuning the regularization parameters Λ in order to minimize the ordinary risk (10) defined above. Yet in practice, the optimal regularization parameters $\Lambda^\dagger$ satisfying

$$\Lambda^\dagger \in \underset{\Lambda \in \mathbb{R}^L}{\text{Argmin}} R[\hat{\mathbf{x}}](\Lambda) \quad (12)$$

is inaccessible. In the following, we propose a detailed procedure to closely approach $\Lambda^\dagger$, by minimizing a Generalized Stein Unbiased Risk Estimator approximating $R[\hat{\mathbf{x}}](\Lambda)$.

2.3 Generalized Stein Unbiased Risk Estimator

The risk defined in (10) depends explicitly on ground truth $\bar{\mathbf{x}}$ and hence is inaccessible. Stein proposed an unbiased estimator of this risk, known as Stein Unbiased Risk Estimator (SURE) in the case of i.i.d. Gaussian noise, recalled in Equation (6). This estimator was then extended to very general noise distributions (see e.g. [30] for Exponential Families, including Gaussian densities). In particular, when the noise ζ is Gaussian, with possible non-trivial covariance matrix, Theorem 1 provides a generalization of Stein's original estimator, which constitutes the starting point of this work.

Stein's approach for risk estimation crucially relies the following hypothesis on estimator $\hat{\mathbf{x}}(\mathbf{y}; \Lambda)$:

Assumption 3 (Regularity and integrability). The estimator $\hat{\mathbf{x}}(\mathbf{y}; \Lambda)$ is continuous and weakly differentiable with respect to observations $\mathbf{y}$. Moreover, the quantities $\langle \mathbf{A}^* \Pi \hat{\mathbf{x}}(\mathbf{y}; \Lambda), \zeta \rangle$ and $\partial_{\mathbf{y}} \hat{\mathbf{x}}(\mathbf{y}; \Lambda)$ are integrable against the Gaussian density:

$$\frac{1}{\sqrt{(2\pi)^P |\det \mathcal{S}|}} \exp \left(-\frac{\|\zeta\|_{\mathcal{S}^{-1}}^2}{2} \right) d\zeta.$$

Theorem 1. Consider Model (1), together with Assumptions 1 (Gaussianity), 2 (Full-rank), 3 (Integrability), and linear operator $\mathbf{A}$ defined in (9). Then generalized Stein's lemma applies, and leads to

$$R[\hat{\mathbf{x}}](\Lambda) = \mathbb{E}_{\zeta} \left[\|\mathbf{A} (\Phi \hat{\mathbf{x}}(\mathbf{y}; \Lambda) - \mathbf{y})\|_2^2 + 2\text{Tr}(\mathcal{S} \mathbf{A}^* \Pi \partial_{\mathbf{y}} \hat{\mathbf{x}}(\mathbf{y}; \Lambda)) - \text{Tr}(\mathcal{S} \mathbf{A} \mathbf{A}^*) \right], \quad (13)$$

the quantity in the brackets being the so-called Generalized Stein Unbiased Risk Estimator.

Proof. A detailed proof is provided in Appendix A. $\square$

Remark 3. Interestingly, when considering the squared Mahalanobis distance in the definition of the risk (11), Stein Unbiased Risk Estimator has the same

global structure, yet, instead of involving the covariance matrix $\mathbf{S}$ it involves its inverse writing

$$\tilde{R}[\hat{\mathbf{x}}](\mathbf{\Lambda}) = \mathbb{E}_{\zeta} \left[\|\mathbf{A}(\Phi\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}) - \mathbf{y})\|_{\mathbf{S}^{-1}}^2 + 2\text{Tr}(\mathbf{A}^* \mathbf{\Pi} \partial_{\mathbf{y}} \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})) - \text{Tr}(\mathbf{A} \mathbf{A}^*) \right]. \quad (14)$$

2.4 Finite Difference Monte Carlo SURE

In the proposed SURE expression (13), the quantity $\text{Tr}(\mathbf{S} \mathbf{A}^* \mathbf{\Pi} \partial_{\mathbf{y}} \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}))$ appearing in (13), called the *degrees of freedom*, concentrates the major difficulties in computing Stein's estimator in data processing problems, as evidenced by the prolific literature addressing this issue in the case $\mathbf{S} \propto \mathbf{I}_P$ [28, 40, 65, 67]. Indeed, it involves the product of the $P \times P$ matrix $\mathbf{S} \mathbf{A}^* \mathbf{\Pi}$ with the $P \times P$ Jacobian matrix $\partial_{\mathbf{y}}(\hat{\mathbf{x}}(\mathbf{y}, \mathbf{\Lambda}))$. Not only the product of two $P \times P$ matrices might be extremely costly in computational efforts but also the Jacobian matrix, because of its large size, $P \gg 1$, might also be very demanding to compute (or even to estimate). Two-step Finite Difference Monte Carlo strategy together with Assumption 4 presented below, enable to overcome these difficulties and to build a usable Stein Unbiased Risk Estimator, denoted $\hat{R}_{\nu, \epsilon}(\mathbf{y}; \mathbf{\Lambda} | \mathbf{S})$, defined in Equation (17).

Assumption 4 (Lipschitzianity w.r.t. observations). Let $\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ an estimator of $\bar{\mathbf{x}}$, depending on observations $\mathbf{y}$, and parametrized by $\mathbf{\Lambda}$.

(i) The mapping $\mathbf{y} \mapsto \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ is uniformly L_1 -Lipschitz .

(ii) $\forall \mathbf{\Lambda} \in \mathbb{R}^L$, $\hat{\mathbf{x}}(\mathbf{0}_P; \mathbf{\Lambda}) = \mathbf{0}_N$, with $\mathbf{0}_N$ (resp. $\mathbf{0}_P$) the null vector of $\mathcal{H}$ (resp. $\mathcal{G}$).

Step 1. Trace estimation via Monte Carlo:

In the way to practical degrees of freedom estimation, the first step is to remark that it far less costly to compute the product of the $P \times P$ matrix $\mathbf{S} \mathbf{A}^* \mathbf{\Pi}$ with $\partial_{\mathbf{y}}(\hat{\mathbf{x}}(\mathbf{y}, \mathbf{\Lambda}))[\epsilon] \in \mathbb{R}^P$, the Jacobian matrix applied on a vector $\epsilon \in \mathbb{R}^P$. Further, straightforward computation shows that if $\epsilon \in \mathbb{R}^P$ is a normalized random variable $\epsilon \sim \mathcal{N}(\mathbf{0}_P, \mathbf{I}_P)$, and $\mathbf{M} \in \mathbb{R}^{P \times P}$ any matrix, then

$$\text{Tr}(\mathbf{M}) = \mathbb{E}_{\epsilon} \langle \mathbf{M} \epsilon, \epsilon \rangle.$$

Thus, following the suggestion of [24, 34, 57], if one has access to $\partial_{\mathbf{y}}(\hat{\mathbf{x}}(\mathbf{y}, \mathbf{\Lambda}))[\epsilon]$, then, since $\mathbf{S}$ is a covariance matrix and hence is symmetric,

$$\begin{aligned} \text{Tr}(\mathbf{S} \mathbf{A}^* \mathbf{\Pi} \partial_{\mathbf{y}} \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})) &= \mathbb{E}_{\epsilon} \langle \mathbf{S} \mathbf{A}^* \mathbf{\Pi} \partial_{\mathbf{y}} \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})[\epsilon], \epsilon \rangle \\ &= \mathbb{E}_{\epsilon} \langle \mathbf{A}^* \mathbf{\Pi} \partial_{\mathbf{y}} \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})[\epsilon], \mathbf{S} \epsilon \rangle, \end{aligned} \quad (15)$$

and $\langle \mathbf{A}^* \mathbf{\Pi} \partial_{\mathbf{y}} \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})[\epsilon], \mathbf{S} \epsilon \rangle$ provides an estimator of degrees of freedom.

Step 2. First-order derivative estimation with Finite Differences:

Second step consists in tackling the problem of estimating $\partial_{\mathbf{y}}(\hat{\mathbf{x}}(\mathbf{y}, \mathbf{\Lambda}))[\epsilon]$ when no direct access to the Jacobian $\partial_{\mathbf{y}}(\hat{\mathbf{x}}(\mathbf{y}, \mathbf{\Lambda}))$ is possible. In this case, the derivative can be estimated using the normalized random variable ϵ and a step $\nu > 0$

making use of Taylor expansion

$$\begin{aligned} \hat{\mathbf{x}}(\mathbf{y} + \nu\boldsymbol{\varepsilon}; \boldsymbol{\Lambda}) - \hat{\mathbf{x}}(\mathbf{y}; \boldsymbol{\Lambda}) &\underset{\nu \rightarrow 0}{\simeq} \partial_{\mathbf{y}}(\hat{\mathbf{x}}(\mathbf{y}; \boldsymbol{\Lambda})) [\nu\boldsymbol{\varepsilon}] \\ \iff \partial_{\mathbf{y}}\hat{\mathbf{x}}(\mathbf{y}; \boldsymbol{\Lambda}) [\boldsymbol{\varepsilon}] &= \lim_{\nu \rightarrow 0} \frac{1}{\nu} (\hat{\mathbf{x}}(\mathbf{y} + \nu\boldsymbol{\varepsilon}; \boldsymbol{\Lambda}) - \hat{\mathbf{x}}(\mathbf{y}; \boldsymbol{\Lambda})). \end{aligned}$$

It follows

$$\begin{aligned} \text{Tr}(\mathbf{S}\mathbf{A}^*\boldsymbol{\Pi}\partial_{\mathbf{y}}\hat{\mathbf{x}}(\mathbf{y}; \boldsymbol{\Lambda})) &= \mathbb{E}_{\boldsymbol{\varepsilon}} \lim_{\nu \rightarrow 0} \frac{1}{\nu} \langle \mathbf{S}\mathbf{A}^*\boldsymbol{\Pi}(\hat{\mathbf{x}}(\mathbf{y} + \nu\boldsymbol{\varepsilon}; \boldsymbol{\Lambda}) - \hat{\mathbf{x}}(\mathbf{y}; \boldsymbol{\Lambda})), \boldsymbol{\varepsilon} \rangle \\ &= \mathbb{E}_{\boldsymbol{\varepsilon}} \lim_{\nu \rightarrow 0} \frac{1}{\nu} \langle \mathbf{A}^*\boldsymbol{\Pi}(\hat{\mathbf{x}}(\mathbf{y} + \nu\boldsymbol{\varepsilon}; \boldsymbol{\Lambda}) - \hat{\mathbf{x}}(\mathbf{y}; \boldsymbol{\Lambda})), \mathbf{S}\boldsymbol{\varepsilon} \rangle. \end{aligned} \quad (16)$$

Elaborating on Formula (16) and Assumption 4, the following theorem provides an asymptotically unbiased Finite Differences Monte Carlo estimator of the risk, which can be used in a vast variety of estimation problems.

Theorem 2. *Consider the observation Model (1), the operator $\mathbf{A}$ defined in (9) together with Assumptions 1 (Gaussianity), 2 (Full-rank), 3 (Integrability), and 4 (Lipschitzianity w.r.t. $\mathbf{y}$). Generalized Finite Differences Monte Carlo SURE, writing*

$$\begin{aligned} \hat{R}_{\nu, \boldsymbol{\varepsilon}}(\mathbf{y}; \boldsymbol{\Lambda} | \mathbf{S}) &\triangleq \|\mathbf{A}(\boldsymbol{\Phi}\hat{\mathbf{x}}(\mathbf{y}; \boldsymbol{\Lambda}) - \mathbf{y})\|_2^2 \\ &+ \frac{2}{\nu} \langle \mathbf{A}^*\boldsymbol{\Pi}(\hat{\mathbf{x}}(\mathbf{y} + \nu\boldsymbol{\varepsilon}; \boldsymbol{\Lambda}) - \hat{\mathbf{x}}(\mathbf{y}; \boldsymbol{\Lambda})), \mathbf{S}\boldsymbol{\varepsilon} \rangle - \text{Tr}(\mathbf{A}\mathbf{S}\mathbf{A}^*), \end{aligned} \quad (17)$$

is an asymptotically unbiased estimator of the risk $R[\hat{\mathbf{x}}](\boldsymbol{\Lambda})$ as $\nu \rightarrow 0$, meaning that

$$\lim_{\nu \rightarrow 0} \mathbb{E}_{\boldsymbol{\zeta}, \boldsymbol{\varepsilon}} \hat{R}_{\nu, \boldsymbol{\varepsilon}}(\mathbf{y}; \boldsymbol{\Lambda} | \mathbf{S}) = R[\hat{\mathbf{x}}](\boldsymbol{\Lambda}). \quad (18)$$

Proof. The proof of Theorem 2 is postponed to Appendix B. $\square$

Remark 4. The use of Monte Carlo strategy is advocated in [24] so that to reduce the complexity of SURE evaluation, replacing costly $P \times P$ matrices product by products of $P \times P$ matrix by vector of size P . Yet, the product of $\mathbf{S} \in \mathbb{R}^{P \times P}$ with $\boldsymbol{\varepsilon} \in \mathbb{R}^P$, as well as the product of $\mathbf{A}^*\boldsymbol{\Pi} \in \mathbb{R}^{P \times N}$ with $\hat{\mathbf{x}}(\mathbf{y}; \boldsymbol{\Lambda}) \in \mathbb{R}^N$, might still be extremely costly. Hopefully, we will see that, in data processing problems (e.g. for texture segmentation in Section 4), both the covariance matrix $\mathbf{S}$ and linear operator $\mathbf{A}$ (through the degradation $\boldsymbol{\Phi}$) benefit from sufficient sparsity so that the calculations can be handled at a reasonable cost.

3 Stein's Unbiased GrAdient estimator of the Risk (SUGAR)

From the estimator of the risk $\hat{R}_{\nu, \boldsymbol{\varepsilon}}(\mathbf{y}; \boldsymbol{\Lambda} | \mathbf{S})$ provided in previous Section 2.3, basic grid search approach could be performed, in order to estimate the optimal $\boldsymbol{\Lambda}^\dagger$, as defined in (12). Yet, the exploration of a fine grid of $\boldsymbol{\Lambda} \in \mathbb{R}^L$ might be

time consuming if the evaluation of $\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ is costly, which is the case when $\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ is *sequential*, i.e. obtained from an optimization scheme. Moreover, the size of a grid in $\mathbb{R}^L$ with given step size grows algebraically with L . Altogether, this precludes grid search when $L > 2$.

Inspiring from [24], this section addresses this issue in the extended case of correlated noise. We provide in Equation (19) a generalized estimator $\partial_{\mathbf{\Lambda}} \hat{R}_{\nu, \epsilon}(\mathbf{y}; \mathbf{\Lambda} | \mathbf{S}) \in \mathbb{R}^L$ of the gradient of the risk with respect to hyperparameters $\mathbf{\Lambda}$. Further, we demonstrate that the Finite Difference Monte Carlo estimator $\partial_{\mathbf{\Lambda}} \hat{R}_{\nu, \epsilon}(\mathbf{y}; \mathbf{\Lambda} | \mathbf{S})$ is an asymptotically unbiased estimator of the gradient of the risk (10) with respect to $\mathbf{\Lambda}$.

In Algorithm 1, we provide an example of sequential estimator, relying on an accelerated primal-dual scheme, designed to solve (3), with its differentiated counterpart, providing both $\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ and its Jacobian $\partial_{\mathbf{\Lambda}} \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}) \in \mathbb{R}^{N \times L}$ with respect to $\mathbf{\Lambda}$.

Hence, costly grid search can be avoided, the estimation of $\mathbf{\Lambda}^\dagger$ being performed by a quasi-Newton descent, described in Algorithm 3, which minimizes the estimated risk $\hat{R}_{\nu, \epsilon}(\mathbf{y}; \mathbf{\Lambda} | \mathbf{S})$, making use of its gradient $\partial_{\mathbf{\Lambda}} \hat{R}_{\nu, \epsilon}(\mathbf{y}; \mathbf{\Lambda} | \mathbf{S})$.

3.1 Differentiation of Stein Unbiased Risk Estimate

Proposition 1. *Consider the observation Model (1), the operator $\mathbf{A}$ defined in (9) together with Assumptions 1 (Gaussianity), 2 (Full-rank), 3 (Integrability), 4 (Lipschitzianity w.r.t. $\mathbf{y}$), and 5 (Lipschitzianity w.r.t. $\mathbf{\Lambda}$) Assumptions. Then the Finite Difference Monte Carlo SURE $\hat{R}_{\nu, \epsilon}(\mathbf{y}; \mathbf{\Lambda} | \mathbf{S})$, defined in (17), is weakly differentiable with respect to both observations $\mathbf{y}$ and parameters $\mathbf{\Lambda}$, and its gradient with respect to $\mathbf{\Lambda}$, as an element of $\mathbb{R}^L$, is given by*

$$\begin{aligned} \partial_{\mathbf{\Lambda}} \left[\hat{R}_{\nu, \epsilon}(\mathbf{y}; \mathbf{\Lambda} | \mathbf{S}) \right] &\triangleq 2 (\mathbf{A} \Phi \partial_{\mathbf{\Lambda}} \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}))^* \mathbf{A} (\Phi \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}) - \mathbf{y}) \\ &+ \frac{2}{\nu} (\mathbf{A}^* \mathbf{\Pi} (\partial_{\mathbf{\Lambda}} \hat{\mathbf{x}}(\mathbf{y} + \nu \epsilon; \mathbf{\Lambda}) - \partial_{\mathbf{\Lambda}} \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})))^* \mathbf{S} \epsilon, \end{aligned} \quad (19)$$

Proof. The Finite Difference Monte Carlo SURE $\hat{R}_{\nu, \epsilon}(\mathbf{y}; \mathbf{\Lambda} | \mathbf{S})$, defined by Formula (17) is a combination of continuous and weakly differentiable functions with respect to both observations $\mathbf{y}$ and parameters $\mathbf{\Lambda}$, composed with (bounded) linear operators, and thus is continuous and weakly differentiable. Further, the derivation rules apply and lead to the expression of Finite Difference Monte Carlo SUGAR estimator given in Formula (19). $\square$

Assumption 5 (Lipschitzianity w.r.t. hyperparameters). Let $\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ be an estimator of $\bar{\mathbf{x}}$, depending on observations $\mathbf{y}$, and parametrized by $\mathbf{\Lambda}$. The mapping $\mathbf{\Lambda} \mapsto \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ is uniformly L_2 -Lipschitz continuous with constant L_2 being independent of $\mathbf{y}$.

Remark 5. As argued in [24], when the estimator $\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ can be expressed as a (composition of) proximal operator(s) of gauge(s) of compact set(s)¹, Assumption 5 holds. Thus, in the case of (3) when $\mathcal{G} = \mathcal{H}$, $\Phi = \mathbf{I}_{\mathcal{H}}$, and $\|\mathbf{U}_{\mathbf{\Lambda}} \mathbf{x}\|_q^q = \lambda \|\mathbf{x}\|_q^q$, for any $q \geq 1$ the Lipschitzianity w.r.t. $\mathbf{\Lambda}$ is ensured. Moreover, in the case of Tikhonov regularization, i.e. $q = 2$ and $\|\mathbf{U}_{\mathbf{\Lambda}} \mathbf{x}\|_q^q = \lambda \|\mathbf{D} \mathbf{x}\|_2^2$

¹For $\mathcal{C} \subset \mathcal{G}$ a non-empty closed convex set containing $\mathbf{0}_{\mathcal{G}}$, the gauge of $\mathcal{C}$ is defined as $\gamma_{\mathcal{C}}(\mathbf{y}) \triangleq \inf \{\omega > 0 \mid \mathbf{y} \in \omega \mathcal{C}\}$.

in (3), if $\Phi = \mathbf{I}_{\mathcal{H}}$ and $\mathbf{D}^*\mathbf{D}$ is diagonalizable with strictly positive eigenvalues, then Assumption 5 is verified. Apart from these two well-known examples, proving the validity of Assumption 5 in the general case of Penalized Least Square is a difficult problem and is foreseen for future work.

Theorem 3. *Consider the observation Model (1), the operator $\mathbf{A}$ defined in (9) together with Gaussianity 1, Full-rank 2, Integrability 3, Lipschitzianity w.r.t. $\mathbf{y}$ 4, and Lipschitzianity w.r.t. $\mathbf{\Lambda}$ 5 Assumptions. Then generalized Finite Difference Monte Carlo SUGAR, $\partial_{\mathbf{\Lambda}}\widehat{R}_{\nu,\varepsilon}(\mathbf{y}; \mathbf{\Lambda})$ defined in Equation (19), is an asymptotically unbiased estimate of the gradient of the risk as $\nu \rightarrow 0$, that is*

$$\partial_{\mathbf{\Lambda}}R[\widehat{\mathbf{x}}](\mathbf{\Lambda}) = \lim_{\nu \rightarrow 0} \mathbb{E}_{\zeta,\varepsilon} \partial_{\mathbf{\Lambda}}\widehat{R}_{\nu,\varepsilon}(\mathbf{y}; \mathbf{\Lambda}|\mathcal{S}) \quad (20)$$

Proof. The proof of Theorem 3 is postponed to Appendix C. $\square$

Remark 6. Finite Difference Monte Carlo estimator of the gradient of the risk, $\partial_{\mathbf{\Lambda}}\widehat{R}_{\nu,\varepsilon}(\mathbf{y}; \mathbf{\Lambda}|\mathcal{S})$, defined in Equation (19), involves the Jacobian $\partial_{\mathbf{\Lambda}}\widehat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}) \in \mathbb{R}^{N \times L}$ which could be a very large matrix, raising difficulties for practical use. Nevertheless, in most applications, the regularization hyperparameters $\mathbf{\Lambda} \in \mathbb{R}^L$, have a “low” dimensionality $L = \mathcal{O}(1) \ll N$. Thus, it is reasonable to expect that the Jacobian matrix $\partial_{\mathbf{\Lambda}}\widehat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}) \in \mathbb{R}^{N \times L}$ can be stored and manipulated, with similar memory and computational costs than for $\widehat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ (see Section 3.2).

3.2 Sequential estimators and forward iterative differentiation

The evaluation of $\partial_{\mathbf{\Lambda}}\widehat{R}_{\nu,\varepsilon}(\mathbf{y}; \mathbf{\Lambda})$ from Formula (19) requires the Jacobian $\partial_{\mathbf{\Lambda}}\widehat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$. Yet, when no closed-form expression of estimator $\widehat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ is available, computing the gradient $\partial_{\mathbf{\Lambda}}\widehat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ might be a complicated task. A large class of estimators $\widehat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ lacking closed-form expression are those obtained as the limit of iterates as

$$\widehat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}) = \lim_{k \rightarrow \infty} \mathbf{x}^{[k]}(\mathbf{y}; \mathbf{\Lambda}), \quad (21)$$

for instance when $\widehat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ is defined as the solution of a minimization problem, e.g. (3).

In the case when $\widehat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ is a sequential estimator, given an observation $\mathbf{y}$, it is only possible to sample the function $\mathbf{\Lambda} \mapsto \widehat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$, for a *discrete* set of regularization hyperparameters $\{\mathbf{\Lambda}_1, \mathbf{\Lambda}_2, \dots\}$, running the minimization algorithm for each hyperparameters $\mathbf{\Lambda}_1, \mathbf{\Lambda}_2, \dots$. It is a classical fact in signal processing that no robust estimator of the differential can be built from samples of the function, thus more sophisticated tools are needed. Provided some smoothness conditions on the iterations of the minimization algorithm, iterative differentiation strategy [24] gives access to a sequence of Jacobian $\partial_{\mathbf{\Lambda}}\widehat{\mathbf{x}}^{[k]}(\mathbf{y}; \mathbf{\Lambda})$, relying on *chain rule* differentiation presented in Proposition 2.

Considering Problem (3), splitting algorithms [4, 20, 50] are advocated to perform the minimization. We chose the primal-dual scheme proposed in [16], Algorithm 2, taking advantage of closed-form expressions of the proximal op-

erators [58] of both the data-fidelity term and the penalization². Chambolle-Pock algorithm, particularized to (3), is presented in Algorithm 1. Further, Φ being full-rank (Assumption 2), denoting by $\text{Sp}(\Phi^* \Phi)$ the spectrum of $\Phi^* \Phi$, $\gamma = 2 \min \text{Sp}(\Phi^* \Phi)$ is strictly positive. Hence the data-fidelity in (3) term turns out to be γ -strongly convex, and the primal-dual algorithm can be accelerated thank to Step (31) of Algorithm 1, following [16]. The iterative differentiation strategy providing $\partial_{\Lambda} \mathbf{x}^{[k]}(\mathbf{y}; \Lambda)$ is presented in the second part of Algorithm 1. Other iterative differentiation schemes are detailed in [24].

Proposition 2. *Let $\Psi : \mathbb{R}^{N \times L} \rightarrow \mathbb{R}^N$ be a differentiable function of variable $\mathbf{x} \in \mathcal{H}$, differentiably parametrized by $\Lambda \in \mathbb{R}^L$, and $(\mathbf{x}^{[k]})_{k \in \mathbb{N}}$ the sequential estimator defined by iterations of the form*

$$\mathbf{x}^{[k+1]} = \Psi(\mathbf{x}^{[k]}; \Lambda). \quad (22)$$

The gradient of $\mathbf{x}^{[k]}$ with respect to Λ can be computed making use of the chain rule differentiation

$$\partial_{\Lambda} \mathbf{x}^{[k+1]} = \partial_{\Lambda} \left(\Psi(\mathbf{x}^{[k]}; \Lambda) \right) = \partial_{\mathbf{x}} \Psi(\mathbf{x}^{[k]}; \Lambda) [\partial_{\Lambda} \mathbf{x}^{[k]}] + \partial_{\Lambda} \Psi(\mathbf{x}^{[k]}; \Lambda), \quad (23)$$

where $\partial_{\mathbf{x}} \Psi(\mathbf{x}; \Lambda) [\delta]$ denotes the differential of Ψ with respect to variable $\mathbf{x}$ applied on vector δ , and $\partial_{\Lambda} \Psi(\mathbf{x}; \Lambda)$ the gradient of Ψ with respect to Λ . The differentiability of Ψ should be understood in the weak sense.

Remark 7. Two particular cases are often encountered in iterative differentiation (see Algorithm 1):

(i) *Linear operator $\Psi(\mathbf{x}; \Lambda) \triangleq \mathbf{U}_{\Lambda} \mathbf{x}$. Assuming that $(\mathbf{x} \mapsto \mathbf{U}_{\Lambda} \mathbf{x})_{\Lambda}$ is a family of linear operators, with a differentiable parametrization by Λ , the chain rule writes*

$$\partial_{\Lambda} \mathbf{x}^{[k+1]} = \mathbf{U}_{\Lambda} \partial_{\Lambda} \mathbf{x}^{[k]} + (\partial_{\Lambda} \mathbf{U}_{\Lambda}) \mathbf{x}^{[k]}, \quad (24)$$

since the differential of the linear operator $\mathbf{U}_{\Lambda}$ with respect to $\mathbf{x}$ is itself. See (33) and (35), in Algorithm 1 for applications of the *chain rule* with linear operators.

(ii) *Proximal operator $\Psi(\mathbf{x}; \Lambda) \triangleq \text{prox}_{\tau \|\cdot\|_{2,1}}(\mathbf{x})$. The proximal operator being independent of Λ , the chain rule simplifies to*

$$\partial_{\Lambda} \mathbf{x}^{[k+1]} = \partial_{\mathbf{x}} \text{prox}_{\tau \|\cdot\|_{2,1}}(\mathbf{x}^{[k]}) [\partial_{\Lambda} \mathbf{x}^{[k]}] \quad (25)$$

with the differential of the so-called $\ell_2 - \ell_1$ *soft-thresholding* $\text{prox}_{\tau \|\cdot\|_{2,1}}$ with respect to $\mathbf{x} = (x_1, x_2)$, applied on $\delta = (\delta_1, \delta_2)$ having the closed-form expression

$$\partial_{\mathbf{x}} \text{prox}_{\tau \|\cdot\|_{2,1}}(\mathbf{x}) [\delta] = \begin{cases} \mathbf{0} & \text{if } \|\mathbf{x}\|_2 \leq \tau \\ \delta - \frac{\tau}{\|\mathbf{x}\|_2} \left(\delta - \frac{\langle \delta, \mathbf{x} \rangle}{\|\mathbf{x}\|_2^2} \mathbf{x} \right) & \text{else.} \end{cases} \quad (26)$$

See (34) and (36), in Algorithm 1 for applications of the *chain rule* with proximal operators.

²see <http://proximity-operator.net> for numerous proximal operator closed-form expressions

Algorithm 1 Accelerated primal-dual scheme for solving (3) with iterative differentiation with respect to regularization parameters Λ .

Routines: $\hat{\mathbf{x}}(\mathbf{y}; \Lambda) = \text{PD}(\mathbf{y}, \Lambda)$
 $(\hat{\mathbf{x}}(\mathbf{y}; \Lambda), \partial_{\Lambda} \hat{\mathbf{x}}(\mathbf{y}; \Lambda)) = \partial \text{PD}(\mathbf{y}, \Lambda)$

Inputs: Observations $\mathbf{y}$
Regularization hyperparameters Λ
Strong-convexity modulus of data-fidelity term
 $\gamma = 2 \min \text{Sp}(\Phi^* \Phi)$

Initialization: Descent steps $\tau^{[0]} = (\tau_1^{[0]}, \tau_2^{[0]})$ such that $\tau_1^{[0]} \tau_2^{[0]} \|\mathbf{U}_{\Lambda}\|^2 < 1$
Primal, auxiliary and dual variables $\mathbf{x}^{[0]} \in \mathcal{H}$, $\mathbf{w}^{[0]} = \mathbf{x}^{[0]}$,
 $\mathbf{z}^{[0]} \in \mathcal{Q}$

Preliminaries: Jacobian with respect to Λ : $\partial_{\Lambda} \mathbf{x}^{[0]}, \partial_{\Lambda} \mathbf{w}^{[0]}, \partial_{\Lambda} \mathbf{z}^{[0]}$

for $k = 1$ **to** $K_{\max}$ **do**

{Accelerated Primal-Dual}

$$\tilde{\mathbf{z}}^{[k]} = \mathbf{z}^{[k]} + \tau_1^{[k]} \mathbf{U}_{\Lambda} \mathbf{w}^{[k]} \quad (27)$$

$$\mathbf{z}^{[k+1]} = \text{prox}_{\tau_1^{[k]}(\|\cdot\|_q^q)^*}(\tilde{\mathbf{z}}^{[k]}) \quad (28)$$

$$\tilde{\mathbf{x}}^{[k]} = \mathbf{x}^{[k]} - \tau_2^{[k]} \mathbf{U}_{\Lambda}^* \mathbf{z}^{[k+1]} \quad (29)$$

$$\mathbf{x}^{[k+1]} = \text{prox}_{\tau_2^{[k]} \|\mathbf{y} - \Phi \cdot\|_2^2}(\tilde{\mathbf{x}}^{[k]}) \quad (30)$$

$$\theta^{[k]} = 1/\sqrt{1 + 2\gamma\tau_2^{[k]}}, \quad \tau_1^{[k+1]} = \tau_1^{[k]}/\theta^{[k]}, \quad \tau_2^{[k+1]} = \theta^{[k]}\tau_2^{[k]} \quad (31)$$

$$\mathbf{w}^{[k+1]} = \mathbf{x}^{[k]} + \theta^{[k]}(\mathbf{x}^{[k+1]} - \mathbf{x}^{[k]}) \quad (32)$$

{Accelerated Differentiated Primal-Dual}

$$\partial_{\Lambda} \tilde{\mathbf{z}}^{[k]} = \partial_{\Lambda} \mathbf{z}^{[k]} + \tau_1^{[k]} \mathbf{U}_{\Lambda} \partial_{\Lambda} \mathbf{w}^{[k]} + \tau_1^{[k]} \frac{\partial \mathbf{U}_{\Lambda}}{\partial \Lambda} \mathbf{w}^{[k]} \quad (33)$$

$$\partial_{\Lambda} \mathbf{z}^{[k+1]} = \partial_{\tilde{\mathbf{z}}} \text{prox}_{\tau_1^{[k]}(\|\cdot\|_q^q)^*}(\tilde{\mathbf{z}}^{[k]}) \left[\partial_{\Lambda} \tilde{\mathbf{z}}^{[k]} \right] \quad (34)$$

$$\partial_{\Lambda} \tilde{\mathbf{x}}^{[k]} = \partial_{\Lambda} \mathbf{x}^{[k]} - \tau_2^{[k]} \mathbf{U}_{\Lambda}^* \partial_{\Lambda} \mathbf{z}^{[k+1]} - \tau_2^{[k]} \frac{\partial \mathbf{U}_{\Lambda}}{\partial \Lambda} \mathbf{z}^{[k+1]} \quad (35)$$

$$\partial_{\Lambda} \mathbf{x}^{[k+1]} = \partial_{\tilde{\mathbf{x}}} \text{prox}_{\tau_2^{[k]} \|\mathbf{y} - \Phi \cdot\|_2^2}(\tilde{\mathbf{x}}^{[k]}) \left[\partial_{\Lambda} \tilde{\mathbf{x}}^{[k]} \right] \quad (36)$$

$$\partial_{\Lambda} \mathbf{w}^{[k+1]} = \partial_{\Lambda} \mathbf{x}^{[k]} + \theta^{[k]}(\partial_{\Lambda} \mathbf{x}^{[k+1]} - \partial_{\Lambda} \mathbf{x}^{[k]}) \quad (37)$$

end for

Outputs: *Finite-time* solution of Problem (3) $\hat{\mathbf{x}}(\mathbf{y}; \Lambda) \triangleq \hat{\mathbf{x}}^{[K_{\max}]}$
Finite-time Jacobian w.r.t. hyperparameters $\partial_{\Lambda} \hat{\mathbf{x}}(\mathbf{y}; \Lambda) \triangleq \partial_{\Lambda} \hat{\mathbf{x}}^{[K_{\max}]}$

Definition 2 (Generalized SURE and SUGAR for sequential estimators). Let $\hat{\mathbf{x}}(\ell; \Lambda)$ be a sequential estimator in the sense of (21). The associated risk es-

timate $\hat{R}_{\nu,\epsilon}(\mathbf{y}; \mathbf{\Lambda}|\mathcal{S})$ and gradient of the risk estimate $\partial_{\mathbf{\Lambda}} \hat{R}_{\nu,\epsilon}(\mathbf{y}; \mathbf{\Lambda}|\mathcal{S})$ are computed running Algorithm 1 twice: first with input $\mathbf{y}$ (observations), second with input $\mathbf{y} + \nu\epsilon$ (perturbed observations). Then, generalized SURE is computed from Formula (17), and generalized SUGAR from Formula (19). These steps are summarized into routines respectively called “SURE” and “SUGAR”, detailed in Algorithm 2.

3.3 Automatic risk minimization

Theorem 2 provides an asymptotically unbiased estimator of the risk $R[\hat{\mathbf{x}}](\mathbf{\Lambda})$, denoted $\hat{R}_{\nu,\epsilon}(\mathbf{y}; \mathbf{\Lambda}|\mathcal{S})$, based on Finite Difference Monte Carlo strategy. Hence, for sufficiently small Finite Difference step $\nu > 0$, we can expect that the solution $\mathbf{\Lambda}^\dagger$ of Problem (7), minimizing the true risk, is well approximated by the hyperparameters $\hat{\mathbf{\Lambda}}_{\nu,\epsilon}^\dagger$ minimizing the estimated risk

$$\hat{\mathbf{\Lambda}}_{\nu,\epsilon}^\dagger(\mathbf{y}|\mathcal{S}) \in \underset{\mathbf{\Lambda} \in \mathbb{R}^L}{\text{Argmin}} \hat{R}_{\nu,\epsilon}(\mathbf{y}; \mathbf{\Lambda}|\mathcal{S}). \quad (45)$$

Then, since the dimensionality of $\mathbf{\Lambda} \in \mathbb{R}^L$ is “low” enough (see Remark 6), Problem (45) is addressed performing a quasi-Newton descent algorithm, using the estimated gradient of the risk $\partial_{\mathbf{\Lambda}} \hat{R}_{\nu,\epsilon}(\mathbf{y}; \mathbf{\Lambda}|\mathcal{S})$, provided by Theorem 3.

A sketch of quasi-Newton descent, particularized to Problem (45), is detailed in Algorithm 3. It generates a sequence $(\mathbf{\Lambda}^{[t]})_{t \in \mathbb{N}}$ converging toward a minimizer of $\hat{R}_{\nu,\epsilon}(\mathbf{y}; \mathbf{\Lambda}|\mathcal{S})$. This algorithm relies on a gradient descent step (48) involving a descent direction $\mathbf{d}^{[t]}$ obtained from the product of BFGS approximated inverse Hessian matrix $\mathbf{H}^{[t]}$ and the gradient $\partial_{\mathbf{\Lambda}} \hat{R}_{\nu,\epsilon}(\mathbf{y}; \mathbf{\Lambda}|\mathcal{S})$ obtained from SUGAR (see Algorithm 2). The descent step size $\alpha^{[t]}$ is obtained from a line search, derived in (47), which stops when Wolfe conditions are fulfilled [22, 49]. Finally, the approximated inverse Hessian matrix $\mathbf{H}^{[t]}$ is updated according to Definition 3.

Remark 8. The line search, Step (47), is the most time consuming. Indeed, the routines SURE and SUGAR are called for several hyperparameters of the form $\mathbf{\Lambda}^{[t]} + \alpha \mathbf{d}^{[t]}$, each call requiring to run differentiated primal-dual scheme twice.

Definition 3 (BroydenFletcherGoldfarbShanno (BFGS)). Let $\mathbf{d}^{[t]}$ be the descent direction and $\mathbf{u}^{[t]}$ the gradient increment at iteration t , the approximated inverse Hessian matrix $\mathbf{H}^{[t]}$ BFGS update writes

$$\mathbf{H}^{[t+1]} = \left(\mathbf{I}_L - \frac{\mathbf{d}^{[t]} (\mathbf{u}^{[t]})^\top}{(\mathbf{u}^{[t]})^\top \mathbf{d}^{[t]}} \right) \mathbf{H}^{[t]} \left(\mathbf{I}_L - \frac{\mathbf{u}^{[t]} (\mathbf{d}^{[t]})^\top}{(\mathbf{u}^{[t]})^\top \mathbf{d}^{[t]}} \right) + \alpha^{[t]} \frac{\mathbf{d}^{[t]} (\mathbf{d}^{[t]})^\top}{(\mathbf{u}^{[t]})^\top \mathbf{d}^{[t]}}. \quad (52)$$

This step constitutes a routine, named “BFGS”, defined as

$$\mathbf{H}^{[t+1]} \triangleq \text{BFGS}(\mathbf{H}^{[t]}, \mathbf{d}^{[t]}, \mathbf{u}^{[t]}). \quad (53)$$

Algorithm 2 Generalized SURE and SUGAR for sequential $\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$.

Routines: $\hat{R}_{\nu, \varepsilon}(\mathbf{y}; \mathbf{\Lambda} | \mathbf{S}) = \text{SURE}(\mathbf{y}, \mathbf{\Lambda}, \mathbf{S}, \nu, \varepsilon)$
 $\partial_{\mathbf{\Lambda}} \hat{R}_{\nu, \varepsilon}(\mathbf{y}; \mathbf{\Lambda} | \mathbf{S}) = \text{SUGAR}(\mathbf{y}, \mathbf{\Lambda}, \mathbf{S}, \nu, \varepsilon)$

Inputs: Observations $\mathbf{y}$
Regularization hyperparameters $\mathbf{\Lambda}$
Covariance matrix $\mathbf{S}$
Monte Carlo vector $\varepsilon \in \mathbb{R}^P \sim \mathcal{N}(\mathbf{0}_P, \mathbf{I}_P)$
Finite Difference step $\nu > 0$

{Solution of (3) from Algorithm 1}

$$\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}) = \text{PD}(\mathbf{y}, \mathbf{\Lambda}) \quad (38)$$

$$\hat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \mathbf{\Lambda}) = \text{PD}(\mathbf{y} + \nu \varepsilon, \mathbf{\Lambda}) \quad (39)$$

{Finite Difference Monte Carlo SURE (17)}

$$\begin{aligned} \hat{R}_{\nu, \varepsilon}(\mathbf{y}; \mathbf{\Lambda} | \mathbf{S}) &= \|\mathbf{A} (\Phi \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}) - \mathbf{y})\|_2^2 \\ &+ \frac{2}{\nu} \langle \mathbf{A}^* \mathbf{\Pi} (\hat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \mathbf{\Lambda}) - \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})), \mathbf{S} \varepsilon \rangle - \text{Tr}(\mathbf{A} \mathbf{S} \mathbf{A}^*) \end{aligned} \quad (40)$$

{Solution of (3) and its differential w.r.t. $\mathbf{\Lambda}$ from Algorithm 1}

$$(\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}), \partial_{\mathbf{\Lambda}} \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})) = \partial \text{PD}(\mathbf{y}, \mathbf{\Lambda}) \quad (41)$$

$$(\hat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \mathbf{\Lambda}), \partial_{\mathbf{\Lambda}} \hat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \mathbf{\Lambda})) = \partial \text{PD}(\mathbf{y} + \nu \varepsilon, \mathbf{\Lambda}) \quad (42)$$

{Finite Difference Monte Carlo estimators (17) and (19)}

$$\begin{aligned} \hat{R}_{\nu, \varepsilon}(\mathbf{y}; \mathbf{\Lambda} | \mathbf{S}) &= \|\mathbf{A} (\Phi \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}) - \mathbf{y})\|_2^2 \\ &+ \frac{2}{\nu} \langle \mathbf{A}^* \mathbf{\Pi} (\hat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \mathbf{\Lambda}) - \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})), \mathbf{S} \varepsilon \rangle - \text{Tr}(\mathbf{A} \mathbf{S} \mathbf{A}^*) \end{aligned} \quad (43)$$

$$\begin{aligned} \partial_{\mathbf{\Lambda}} \hat{R}_{\nu, \varepsilon}(\mathbf{y}; \mathbf{\Lambda} | \mathbf{S}) &= 2 (\mathbf{A} \Phi \partial_{\mathbf{\Lambda}} \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}))^* \mathbf{A} (\Phi \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}) - \mathbf{y}) \\ &+ \frac{2}{\nu} (\mathbf{A}^* \mathbf{\Pi} (\partial_{\mathbf{\Lambda}} \hat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \mathbf{\Lambda}) - \partial_{\mathbf{\Lambda}} \hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})))^* \mathbf{S} \varepsilon \end{aligned} \quad (44)$$

Output: Risk estimate $\hat{R}_{\nu, \varepsilon}(\mathbf{y}; \mathbf{\Lambda} | \mathbf{S})$
Gradient of the risk estimate $\partial_{\mathbf{\Lambda}} \hat{R}_{\nu, \varepsilon}(\mathbf{y}; \mathbf{\Lambda} | \mathbf{S})$

For detailed discussions on low memory implementations of BFGS, box constraints management, and others algorithmic tricks the interested reader is referred to [12, 22, 49].

Convergence conditions for quasi-Newton algorithms relies on the behavior of

Algorithm 3 Automated selection of hyperparameters minimizing quadratic risk.

Inputs: Observations $\mathbf{y}$

Covariance matrix $\mathcal{S}$

Monte Carlo vector $\boldsymbol{\varepsilon} \in \mathbb{R}^P \sim \mathcal{N}(\mathbf{0}_P, \mathbf{I}_P)$

Finite Difference step $\nu > 0$

Initialization: $\boldsymbol{\Lambda}^{[0]} \in \mathbb{R}^L$, Inverse Hessian $\mathbf{H}^{[0]} \in \mathbb{R}^{L \times L}$,
Gradient $\partial_{\boldsymbol{\Lambda}} \widehat{R}^{[0]} = \text{SUGAR}(\mathbf{y}, \boldsymbol{\Lambda}^{[0]}, \mathcal{S}, \nu, \boldsymbol{\varepsilon})$

for $t = 0$ **to** $T_{\max} - 1$ **do**

{Descent direction from gradient of the risk estimate:}

$$\mathbf{d}^{[t]} = -\mathbf{H}^{[t]} \partial_{\boldsymbol{\Lambda}} \widehat{R}^{[t]} \quad (46)$$

{Line search to find descent step:}

$$\alpha^{[t]} \in \underset{\alpha \in \mathbb{R}}{\text{Argmin}} \widehat{R}(\boldsymbol{\Lambda}^{[t]} + \alpha \mathbf{d}^{[t]}), \quad \text{with } \widehat{R}(\boldsymbol{\Lambda}) = \text{SURE}(\mathbf{y}, \boldsymbol{\Lambda}, \mathcal{S}, \nu, \boldsymbol{\varepsilon}) \quad (47)$$

{Quasi-Newton descent step on $\boldsymbol{\Lambda}$:}

$$\boldsymbol{\Lambda}^{[t+1]} = \boldsymbol{\Lambda}^{[t]} + \alpha^{[t]} \mathbf{d}^{[t]} \quad (48)$$

{Gradient update:}

$$\partial_{\boldsymbol{\Lambda}} \widehat{R}^{[t+1]} = \text{SUGAR}(\mathbf{y}, \boldsymbol{\Lambda}^{[t+1]}, \mathcal{S}, \nu, \boldsymbol{\varepsilon}) \quad (49)$$

{Gradient increment}

$$\mathbf{u}^{[t]} = \partial_{\boldsymbol{\Lambda}} \widehat{R}^{[t+1]} - \partial_{\boldsymbol{\Lambda}} \widehat{R}^{[t]} \quad (50)$$

{BFGS update of inverse Hessian (52):}

$$\mathbf{H}^{[t+1]} = \text{BFGS}(\mathbf{H}^{[t]}, \mathbf{d}^{[t]}, \mathbf{u}^{[t]}) \quad (51)$$

end for

Outputs: *Finite-time* solution of Problem (45) $\widehat{\boldsymbol{\Lambda}}_{\nu, \boldsymbol{\varepsilon}}^{\text{BFGS}}(\mathbf{y}|\mathcal{S}) \triangleq \boldsymbol{\Lambda}^{[T_{\max}]}$
Estimate with automated selection of $\boldsymbol{\Lambda}$ $\widehat{\mathbf{x}}_{\nu, \boldsymbol{\varepsilon}}^{\text{BFGS}}(\mathbf{y}|\mathcal{S}) \triangleq$
PD($\mathbf{y}, \boldsymbol{\Lambda}^{[T_{\max}]}$)

second derivatives of the objective function [49]. Most of the time, when it comes to sequential estimators, one has no information about the twice differentiability of generalized SURE with respect to hyperparameters. Hence, the convergence of Algorithm 3 will be assessed numerically. Further, quasi-Newton algorithms being known to be sensitive to initialization, special attention needs to be paid to the initialization of both hyperparameters $\boldsymbol{\Lambda}$ and approximated inverse Hessian

H (see Section 5.2.4).

Remark 9. Given a parametric estimator $\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$, possibly obtained by another routine than PD, Algorithms 2 and 3 can be used, provided that one has a routine equivalent to ∂PD , computing $\partial_{\mathbf{\Lambda}}\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$. The reader can find other differentiated proximal algorithms in [24].

4 Hyperparameter tuning for texture segmentation

The formalism proposed above for the automated selection of the regularization hyperparameters is now specified to total-variation based texture segmentation. Section 4.1 formulates the texture segmentation problem as the minimization of a convex objective function. Then, in Section 4.2, this segmentation procedure is cast into the general formalism of Sections 2 and 3. The hypothesis needed to apply Theorems 2 and 3 are discussed one by one in the context of texture segmentation. Finally, the practical evaluation of the estimators of the risk $\hat{R}_{\nu,\epsilon}(\ell; \mathbf{\Lambda}|\mathcal{S})$ and of the gradient of the risk $\partial_{\mathbf{\Lambda}}\hat{R}_{\nu,\epsilon}(\ell; \mathbf{\Lambda}|\mathcal{S})$ is discussed in Section 4.3.

4.1 Total-variation based texture segmentation

4.1.1 Piecewise homogeneous fractal texture model

Let $X \in \mathbb{R}^{N_1 \times N_2}$ denote the texture to be segmented, consisting of a real-valued discrete field defined on a grid of pixels $\Omega = \{1, \dots, N_1\} \times \{1, \dots, N_2\}$. Texture X is assumed to be formed as the union of M independent Gaussian textures, existing on a set of disjoint supports,

$$\Omega = \overline{\Omega}_1 \cup \dots \cup \overline{\Omega}_M, \quad \text{with} \quad \overline{\Omega}_m \cap \overline{\Omega}_{m'} = \emptyset \quad \text{if} \quad m \neq m'. \quad (54)$$

Each homogeneous Gaussian texture, defined on Ω_m is characterized by two global fractal features, the scaling (or Hurst) exponent $\overline{H}_m$ and the variance $\overline{\Sigma}_m^2$, that fully control its statistics. Interested readers are referred to e.g., [51] for the detailed definition of Gaussian fractal textures. Figures 1b and 1c propose examples of such piecewise Gaussian fractal textures, with $M = 2$ and mask shown in Figure 1a.

4.1.2 Local regularity and wavelet *leader* coefficients

It was abundantly discussed in the literature (cf. e.g. [48, 54, 70–72]) that textures can be well-analyzed by local fractal features (local regularity and local variance), that can be accurately estimated from wavelet *leader* coefficients, as extensively described and studied in e.g. [56, 72], to which the reader is referred for a detailed presentation.

Let $\chi_{j,\underline{n}}^{(d)}$ denote the coefficients of the undecimated 2D Discrete Wavelet Transform of image X , at octave $j = j_1, \dots, j_2$ and pixel $\underline{n} \in \Omega$, with the 2D-wavelet basis being defined from the 4 combination (hence the orientations $d \in \{0, 1, 2, 3\}$) of 1D wavelet ψ and scaling functions. Interested readers are referred to e.g., [46] for a full definition of the $\chi_{j,\underline{n}}^{(d)}$. Wavelet *leaders*,

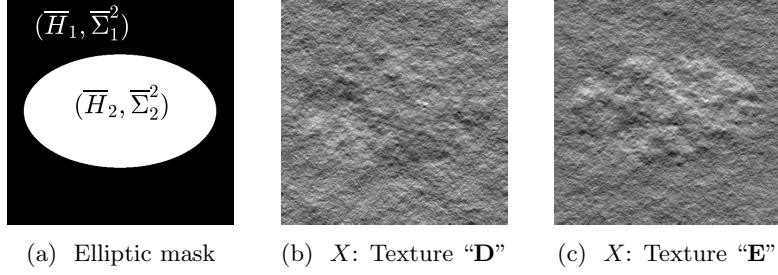


Figure 1: (a) Mask for piecewise textures composed of two regions: “background” (in black) on which the texture is characterized by homogeneous local regularity $\bar{h} \equiv \bar{H}_0$ and local variance $\bar{\sigma}^2 \equiv \bar{\Sigma}_0^2$ and “foreground” (in white) on which the texture is characterized by homogeneous local regularity $\bar{h} \equiv \bar{H}_1$ and local variance $\bar{\sigma}^2 \equiv \bar{\Sigma}_1^2$. (b) and (c) Synthetic piecewise homogeneous textures used for performance assessment, with resolution 256×256 pixels.

$\{\mathcal{L}_{j,\underline{n}}, j = j_1, \dots, j_2, \underline{n} \in \Omega\}$, are further defined as local suprema over a spatial neighborhood and across all finest scales of the $\chi_{j,\underline{n}}^{(d)}$ [72]:

$$\mathcal{L}_{j,\underline{n}} = \sup_{\substack{d = \{1, 2, 3\} \\ \lambda_{j',\underline{n}'} \subset 3\lambda_{j,\underline{n}}}} \left| 2^j \chi_{j',\underline{n}'}^{(d)} \right|, \quad \text{where} \quad \begin{cases} \lambda_{j,\underline{n}} = [\underline{n}, \underline{n} + 2^j[, \\ 3\lambda_{j,\underline{n}} = \bigcup_{\underline{p} \in \{-2^j, 0, 2^j\}^2} \lambda_{j,\underline{n}+\underline{p}}. \end{cases} \quad (55)$$

Local regularity $\bar{h}_{\underline{n}}$ and local variance $\bar{\sigma}_{\underline{n}}^2$ at pixel n can be defined via the local power law behavior of the wavelet leaders across scales [71, 72]:

$$\mathcal{L}_{j,\underline{n}} = \bar{\sigma}_{\underline{n}} 2^{j\bar{h}_{\underline{n}}} \beta_{j,\underline{n}}, \quad \text{as } 2^j \rightarrow 0, \quad (56)$$

where $\beta_{j,\underline{n}}$ can be well approximated for large classes of textures [70] as log-normal random variables, with log-mean $\mu = 0$. For piecewise fractal textures X described in Section 4.1.1, local regularity $\bar{h} \in \mathbb{R}^{N_1 \times N_2}$ and local variance $\bar{\sigma}^2 \in \mathbb{R}^{N_1 \times N_2}$ maps are piecewise constant, reflecting the global scaling exponent H and variance Σ^2 of the homogeneous textures as:

$$(\forall m \in \{1, \dots, M\}) \quad (\forall \underline{n} \in \bar{\Omega}_m) \quad \bar{h}_{\underline{n}} \equiv \bar{H}_m \text{ and } \bar{\sigma}_{\underline{n}}^2 \equiv \bar{\Sigma}_m^2 F(\bar{H}_m, \psi), \quad (57)$$

with $F(\bar{H}_m, \psi)$ a deterministic function studied in [68] and not of interest here. Taking the logarithm of Equation (56) leads to the following linear formulation

$$\ell_{j,\underline{n}} = \bar{v}_{\underline{n}} + j\bar{h}_{\underline{n}} + \zeta_{j,\underline{n}}, \quad \text{as } 2^j \rightarrow 0 \quad (58)$$

with log-leaders $\ell_{j,\underline{n}} = \log_2(\mathcal{L}_{j,\underline{n}})$, log-variance $\bar{v}_{\underline{n}} = \log_2 \bar{\sigma}_{\underline{n}}$ and zero-mean Gaussian noise $\zeta_{j,\underline{n}} = \log_2(\beta_{j,\underline{n}})$. In the following, the *leader* coefficients at scale 2^j are denoted $\ell_j \in \mathbb{R}^{N_1 N_2}$, and the complete collection of *leaders* is stored in $\ell \in \mathbb{R}^{J N_1 N_2}$.

4.1.3 Total variation regularization and iterative thresholding

The linear regression estimator inspired by (58)

$$\begin{pmatrix} \hat{\mathbf{h}}_{\text{LR}}(\boldsymbol{\ell}) \\ \hat{\mathbf{v}}_{\text{LR}}(\boldsymbol{\ell}) \end{pmatrix} = \underset{\begin{pmatrix} \mathbf{h} \\ \mathbf{v} \end{pmatrix} \in \mathbb{R}^{2N_1N_2}}{\text{argmin}} \sum_{j=j_1}^{j_2} \|j\mathbf{h} + \mathbf{v} - \boldsymbol{\ell}_j\|_2^2 \quad (59)$$

achieves poor performance in estimating piecewise constant local regularity and local power, hence precluding an accurate segmentation of the piecewise homogeneous textures. Thus, a functional for *joint* attribute estimation and segmentation was proposed in [51], leading to the following Penalized Least Squares (3):

$$\begin{pmatrix} \hat{\mathbf{h}}(\boldsymbol{\ell}; \boldsymbol{\Lambda}) \\ \hat{\mathbf{v}}(\boldsymbol{\ell}; \boldsymbol{\Lambda}) \end{pmatrix} \in \underset{\begin{pmatrix} \mathbf{h} \\ \mathbf{v} \end{pmatrix} \in \mathbb{R}^{2N_1N_2}}{\text{Argmin}} \sum_{j=j_1}^{j_2} \|j\mathbf{h} + \mathbf{v} - \boldsymbol{\ell}_j\|_2^2 + \lambda_h \text{TV}(\mathbf{h}) + \lambda_v \text{TV}(\mathbf{v}), \quad (60)$$

where TV stands for the well-known isotropic Total Variation, defined as a mixed $\ell_{2,1}$ -norm composed with spatial gradient operators

$$\text{TV}(\mathbf{h}) = \sum_{\underline{n} \in \Omega} \sqrt{(\mathbf{D}_1 \mathbf{h})_{\underline{n}}^2 + (\mathbf{D}_2 \mathbf{h})_{\underline{n}}^2} = \sum_{\underline{n} \in \Omega} \|(\mathbf{D} \mathbf{h})_{\underline{n}}\|_2, \quad (61)$$

where $\mathbf{D}_1 : \mathbb{R}^{N_1N_2} \rightarrow \mathbb{R}^{N_1N_2}$ (resp. $\mathbf{D}_2 : \mathbb{R}^{N_1N_2} \rightarrow \mathbb{R}^{N_1N_2}$) stand for the discrete spatial horizontal (resp. vertical) gradient operator. This TV-penalized least square estimator is designed to favor piecewise constancy of the estimates $\hat{\mathbf{h}}$ and $\hat{\mathbf{v}}$, making used of ℓ_1 -norm, i.e. $q = 1$ in (3).

Finally, following [14, 15], the estimate $\hat{\mathbf{h}}(\boldsymbol{\ell}; \boldsymbol{\Lambda})$ is *thresholded* to yield a posterior *piecewise constant* map of local regularity $T\hat{\mathbf{h}}(\boldsymbol{\ell}; \boldsymbol{\Lambda})$, taking *exactly* M different values $\hat{H}_1(\boldsymbol{\ell}; \boldsymbol{\Lambda}), \dots, \hat{H}_M(\boldsymbol{\ell}; \boldsymbol{\Lambda})$. The resulting segmentation

$$\Omega = \hat{\Omega}_1(\boldsymbol{\ell}; \boldsymbol{\Lambda}) \cup \dots \cup \hat{\Omega}_M(\boldsymbol{\ell}; \boldsymbol{\Lambda}) \quad (62)$$

is deduced from $T\hat{\mathbf{h}}(\boldsymbol{\ell}; \boldsymbol{\Lambda})$, defining

$$(\forall m \in \{1, \dots, M\}), \hat{\Omega}_m(\boldsymbol{\ell}; \boldsymbol{\Lambda}) = \left\{ \underline{n} \in \Omega \mid \left(T\hat{\mathbf{h}}(\boldsymbol{\ell}; \boldsymbol{\Lambda}) \right)_{\underline{n}} \equiv \hat{H}_m(\boldsymbol{\ell}; \boldsymbol{\Lambda}) \right\}. \quad (63)$$

This is illustrated in Figure 2, for a two-region synthetic texture with ground truth piecewise constant local regularity $\bar{\mathbf{h}}$ in Figure 2a.

4.2 Reformulation in term of Model (1)

4.2.1 Observation $\mathbf{y}$

To cast the log-linear behavior (58) into the general model (1), vectorized quantities for $\bar{\mathbf{v}}$, $\bar{\mathbf{h}}$ and $\boldsymbol{\ell}$ are used. The $N_1 \times N_2$ maps $\bar{\mathbf{h}}$ and $\bar{\mathbf{v}}$ are reshaped into vectors $\bar{\mathbf{x}} \in \mathbb{R}^N$, with $N = 2N_1N_2$, ordering the pixels in the lexicographic order. The log-leaders $\boldsymbol{\ell} = (\boldsymbol{\ell}_j)_{j_1 \leq j \leq j_2}$, composed of $J \triangleq j_2 - j_1 + 1$ octaves of resolution $N_1 \times N_2$ are vectorized, octaves by octaves, with lexical ordering of pixels, $\boldsymbol{\ell} \in \mathbb{R}^P$, with $P = JN_1N_2$.

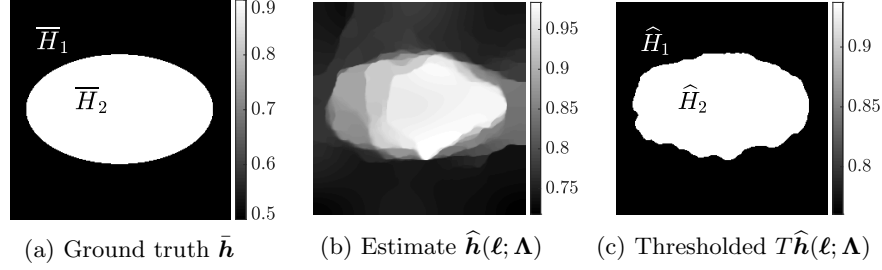


Figure 2: Example of the thresholding of $\hat{\mathbf{h}}(\ell; \Lambda)$ to obtain a two-region segmentation where ℓ denotes the wavelet *learders* associated with the Texture X : “E” displayed in Figure 1c

Equation (58) can then be cast into general model (1) as:

$$\text{Observations} \quad \mathbf{y} = \ell \in \mathbb{R}^P, \quad P = JN_1N_2 \quad (64)$$

$$\text{Ground truth} \quad \bar{\mathbf{x}} = \begin{pmatrix} \bar{\mathbf{h}} \\ \bar{\mathbf{v}} \end{pmatrix} \in \mathbb{R}^N, \quad N = 2N_1N_2 \quad (65)$$

$$\text{Linear degradation} \quad \Phi : \begin{cases} \mathbb{R}^N \rightarrow \mathbb{R}^P \\ \begin{pmatrix} \bar{\mathbf{h}} \\ \bar{\mathbf{v}} \end{pmatrix} \mapsto (j\bar{\mathbf{h}} + \bar{\mathbf{v}})_{j_1 \leq j \leq j_2} \end{cases} \quad (66)$$

4.2.2 Full-rank operator Φ

Proposition 3 asserts that $\Phi^*\Phi$ is invertible (Assumption 2).

Proposition 3. *The linear operator Φ defined in (66) is bounded and its adjoint writes*

$$\Phi^* : \begin{cases} \mathbb{R}^P \rightarrow \mathbb{R}^N \\ (\ell_j)_{j_1 \leq j \leq j_2} \mapsto \begin{pmatrix} \sum_{j=j_1}^{j_2} j\ell_j \\ \sum_{j=j_1}^{j_2} \ell_j \end{pmatrix} \end{cases} \quad (67)$$

Further, Φ is full rank, and the following inversion formula holds

$$(\Phi^*\Phi)^{-1} = \frac{1}{F_2F_0 - F_1^2} \begin{pmatrix} F_0\mathbf{I}_{N/2} & -F_1\mathbf{I}_{N/2} \\ -F_1\mathbf{I}_{N/2} & F_2\mathbf{I}_{N/2} \end{pmatrix}, \quad F_\alpha \triangleq \sum_{j=j_1}^{j_2} j^\alpha, \quad \alpha \in \{0, 1, 2\}. \quad (68)$$

Proof. Formula (67) is obtained from straightforward computations. Then, combining (66) and (67), leads to

$$\Phi^*\Phi = \begin{pmatrix} F_2\mathbf{I}_{N/2} & F_1\mathbf{I}_{N/2} \\ F_1\mathbf{I}_{N/2} & F_0\mathbf{I}_{N/2} \end{pmatrix} \quad (69)$$

which is finally inverted using the 2×2 cofactor matrix formula. $\square$

4.2.3 Projection operator

Performing texture segmentation the discriminant attribute is the local regularity $\mathbf{h}$, while local power $\mathbf{v}$ is an auxiliary feature. Hence the *projected* quadratic risk (10) customized to texture segmentation reads:

$$R[\hat{\mathbf{h}}](\Lambda) \triangleq \mathbb{E}_{\zeta} \left\| \hat{\mathbf{h}}(\ell; \Lambda) - \bar{\mathbf{h}} \right\|_2^2, \quad (70)$$

with $\hat{\mathbf{h}}(\ell; \Lambda)$ defined in (60).

Then, the particularized projection operator in Definition 1 takes the matrix form

$$\Pi \triangleq \begin{pmatrix} \mathbf{I}_{N/2} & \mathbf{Z}_{N/2} \\ \mathbf{Z}_{N/2} & \mathbf{Z}_{N/2} \end{pmatrix} \quad \text{so that} \quad \Pi \begin{pmatrix} \bar{\mathbf{h}} \\ \mathbf{0}_{N/2} \end{pmatrix} = \begin{pmatrix} \bar{\mathbf{h}} \\ \mathbf{0}_{N/2} \end{pmatrix} \quad (71)$$

where $\mathbf{I}_{N/2}$ (resp. $\mathbf{Z}_{N/2}$) denotes the identity (resp. null) matrix of size $N/2 \times N/2$ and $\mathbf{0}_{N/2}$ the null vector of $\mathbb{R}^{N/2}$.

4.2.4 Regularity of the estimates

Proposition 4. *Problem (60) has a unique solution $\begin{pmatrix} \hat{\mathbf{h}}(\ell; \Lambda) \\ \hat{\mathbf{v}}(\ell; \Lambda) \end{pmatrix}$.*

This solution is continuous and weakly differentiable w.r.t. ℓ and integrable against the Gaussian probability density function (Assumption 3). Further, both $\hat{\mathbf{h}}(\ell; \Lambda)$ and $\hat{\mathbf{v}}(\ell; \Lambda)$ are uniformly L_1 -Lipschitz w.r.t. ℓ (Assumption 4).

Proof. As shown in [51], the objective function

$$(\mathbf{h}, \mathbf{v}) \mapsto \underbrace{\sum_{j=j_1}^{j_2} \|\mathbf{j}\mathbf{h} + \mathbf{v} - \ell_j\|_2^2}_{\left\| \Phi \begin{pmatrix} \mathbf{h} \\ \mathbf{v} \end{pmatrix} - \ell \right\|_2^2} + \lambda_h \text{TV}(\mathbf{v}) + \lambda_v \text{TV}(\mathbf{v}) \quad (72)$$

is convex, being the sum of convex terms. Further, computing the eigenvalues of $\Phi^* \Phi$ shows that the least squares data fidelity term is γ -strongly convex, with

$$\gamma = 2 \min \text{Sp}(\Phi^* \Phi) > 0 \quad (73)$$

where $\text{Sp}(\Phi^* \Phi)$ stand for the spectrum of the (bounded) linear operator $\Phi^* \Phi$. Hence, the objective function (72) has a unique minimum, being the unique solution of Problem (60), as mentioned in Remark 1.

Further, (60) falls under the general formulation of Penalized Least Squares (3), which can be written

$$\hat{\mathbf{x}}(\mathbf{y}; \Lambda) = \underset{\mathbf{x} \in \mathcal{H}}{\text{argmin}} \|\mathbf{y} - \Phi \mathbf{x}\|_{\mathbf{W}}^2 + \mathcal{J}_{\Lambda}(\mathbf{x}), \quad (74)$$

where $\mathcal{J}_{\Lambda}(\mathbf{x}) = \|\mathbf{U}_{\Lambda} \mathbf{x}\|_1$ is built from a linear operator $\mathbf{U}_{\Lambda}$ depending on regularization parameters λ_h and λ_v as

$$\Lambda = \begin{pmatrix} \lambda_h \\ \lambda_v \end{pmatrix} \in \mathbb{R}_+^2, \quad \text{and} \quad \mathbf{U}_{\Lambda} = \begin{cases} \mathbb{R}^{2N_1 N_2} \rightarrow \mathbb{R}^{2N_1 N_2} \times \mathbb{R}^{2N_1 N_2} \\ \begin{pmatrix} \mathbf{h} \\ \mathbf{v} \end{pmatrix} \mapsto \left(\begin{pmatrix} \lambda_h \mathbf{D}_1 \mathbf{h} \\ \lambda_v \mathbf{D}_1 \mathbf{v} \end{pmatrix}, \begin{pmatrix} \lambda_h \mathbf{D}_2 \mathbf{h} \\ \lambda_v \mathbf{D}_2 \mathbf{v} \end{pmatrix} \right) \end{cases}. \quad (75)$$

$\mathcal{J}_\Lambda$ is convex, proper and lower semicontinuous, then, following [67], (74) can be rewritten as a constrained optimization problem

$$(\hat{\mathbf{x}}(\mathbf{y}; \Lambda), \hat{\mathbf{z}}(\mathbf{y}; \Lambda)) = \underset{\mathbf{x} \in \mathcal{H}, \mathbf{z} \in \mathcal{G}}{\operatorname{argmin}} \|\mathbf{y} - \mathbf{z}\|_{\mathcal{W}}^2 + \mathcal{J}_\Lambda(\mathbf{x}), \text{ such that } \mathbf{z} = \Phi \mathbf{x} \quad (76)$$

$$\iff \hat{\mathbf{z}}(\mathbf{y}; \Lambda) = \underset{\mathbf{z} \in \mathcal{G}}{\operatorname{argmin}} \|\mathbf{y} - \mathbf{z}\|_{\mathcal{W}}^2 + (\Phi \mathcal{J}_\Lambda)(\mathbf{z}) \quad (77)$$

$$\iff \hat{\mathbf{z}}(\mathbf{y}; \Lambda) = \operatorname{prox}_{(1/2)\Phi \mathcal{J}_\Lambda}(\mathbf{y}) \quad (78)$$

where

$$(\Phi \mathcal{J}_\Lambda)(\mathbf{z}) \triangleq \min_{\{\mathbf{x} | \Phi \mathbf{x} = \mathbf{z}\}} \mathcal{J}_\Lambda(\mathbf{x}) \quad (79)$$

denotes the pre-image of $\mathcal{J}_\Lambda$ under Φ , which is as well convex, proper and lower semicontinuous.

Then, from (78), the estimator $\hat{\mathbf{z}}(\mathbf{y}; \Lambda)$ is non expansive, i.e. 1-Lipschitz, because the proximal operators share that same property. Moreover, from (76), $\hat{\mathbf{z}}(\mathbf{y}; \Lambda) = \Phi \hat{\mathbf{x}}(\mathbf{y}; \Lambda)$, and since Φ is full-rank according to Proposition 3

$$\hat{\mathbf{x}}(\mathbf{y}; \Lambda) = (\Phi^* \Phi)^{-1} \Phi^* \hat{\mathbf{z}}(\mathbf{y}; \Lambda). \quad (80)$$

Φ being bounded, we conclude that the estimator $\hat{\mathbf{x}}(\mathbf{y}; \Lambda)$ is uniformly L_1 -Lipschitz, with $L_1 = \|\Phi\|^{-1}$ justifying Assumption 4, (i).

Being uniformly L_1 -Lipschitz, $\hat{\mathbf{x}}(\mathbf{y}; \Lambda)$ is continuous and weakly-differentiable (see Theorem 5 of Section 4.2.3 in [32]). As a consequence, both $\langle \mathbf{A}^* \Pi \hat{\mathbf{x}}(\mathbf{y}; \Lambda), \zeta \rangle$ and $\partial_\Lambda \hat{\mathbf{x}}(\mathbf{y}; \Lambda)$ are integrable against the Gaussian density and Assumption 3 holds.

Finally, setting $\mathbf{y} = \mathbf{0}_P$, for any $\Lambda \in \mathbb{R}^L$, $\hat{\mathbf{x}}(\mathbf{0}_P; \Lambda) = \mathbf{0}_N$ reaches the minimum. The solution being unique from Proposition 3, $\mathbf{0}_N$ is the unique solution and Assumption 4, (ii) is verified.

Further, it is reasonable to expect that the uniform Lipschitzianity with respect to hyperparameters results of Remark 5, extend to the estimator $\hat{\mathbf{h}}(\ell; \Lambda)$, defined in (60). Yet, to the best of our knowledge, no direct proof that Lipschitzianity Assumption 5 holds for general Penalized Least Squares exists. This issue is a scientific question in itself and will be addressed in future work. $\square$

4.3 Practical computation of $\hat{R}_{\nu, \varepsilon}$ and $\partial_\Lambda \hat{R}_{\nu, \varepsilon}$

This section addresses all technical issues encountered in running Algorithm 2, in the context of texture segmentation described above.

4.3.1 Covariance structure of the observations

The additive noise $\zeta_{j, \underline{n}}$ appearing in Equation (58) being Gaussian, Gaussianity Assumption 1 holds. The covariance matrix $\mathcal{S}$ of noise ζ reads

$$\mathcal{S}_{j, \underline{n}}^{j', \underline{n}'} \triangleq \mathbb{E} \zeta_{j, \underline{n}} \zeta_{j', \underline{n}'} = \mathcal{C}_j^{j'} \Xi_j^{j'}(\underline{n} - \underline{n}'), \quad (81)$$

where

$$\mathcal{C}_j^{j'} \triangleq \mathbb{E} \zeta_{j, \underline{n}} \zeta_{j', \underline{n}}, \quad \mathcal{C}_j^{j'} \text{ independent of } \underline{n} \quad (82)$$

quantifies the inter-scale covariance, and $\Xi_j^{j'}$ encapsulate the stationary spatial correlations, with correlation length proportional to $\max(2^j, 2^{j'})$.

4.3.2 Matrix product $\mathcal{S}\varepsilon$

Following Remark 4, in general, the direct product $\mathcal{S}\varepsilon$ required for the practical evaluation of Finite Difference Monte Carlo SURE (17) is intractable because of the large size of matrix $\mathcal{S}$. Yet, in the case of log-*leaders*, the spatial correlations presenting the Toeplitz structure (81), the product $\mathcal{S}\varepsilon$ can be computed efficiently, making (17) usable in practice.

Indeed, given $\varepsilon \in \mathbb{R}^{JN_1N_2} = (\varepsilon_j)_{j=j_1}^{j_2} J$, with $\varepsilon_j \in \mathbb{R}^{N_1N_2}$,

$$(\mathcal{S}\varepsilon)_{j,\underline{n}} = \sum_{j'=j_1}^{j_2} \sum_{\underline{n}' \in \Omega} \mathcal{S}_{j,\underline{n}}^{j',\underline{n}'} \varepsilon_{j',\underline{n}'} = \sum_{j'=j_1}^{j_2} \mathcal{C}_j^{j'} \sum_{\underline{n}' \in \Omega} \Xi_j^{j'}(\underline{n} - \underline{n}') \varepsilon_{j',\underline{n}'} = \sum_{j'=j_1}^{j_2} \mathcal{C}_j^{j'} \Xi_{j,j'}^{j'} * \varepsilon_j. \quad (83)$$

which is the sum of J convolution products, denoted $*$, of high dimensional vector $\varepsilon \in \mathbb{R}^{N_1N_2}$ with “low dimensional” finite support window $\mathcal{C}_j^{j'} \Xi_{j,j'}^{j'}$. Hence evaluating $\mathcal{S}\varepsilon$ appears to be far less costly than a general product of matrix of size $P \times P$ by a vector of size P .

4.3.3 Operator $\mathbf{A}$

In the same vein, the matrices Φ^* (67), $(\Phi^* \Phi)^{-1}$ (68) and Π (71) turn out to be very sparse, since they act independently on each pixel. Thus, the same sparse (pixel-wise) structure follows for

$$\mathbf{A} = \Pi (\Phi^* \Phi)^{-1} \Phi^* = \frac{1}{F_0 F_2 - F_1^2} \begin{pmatrix} (F_0 + j_1 F_1) \mathbf{I}_{N/2} & \cdots & (F_0 + j_2 F_1) \mathbf{I}_{N/2} \\ \mathbf{Z}_{N/2} & \cdots & \mathbf{Z}_{N/2} \end{pmatrix}. \quad (84)$$

Hence the products $\mathbf{A}^* \Pi \hat{x}(\mathbf{y}; \Lambda)$ and $\mathbf{A}^* \Pi \partial_\Lambda \hat{x}(\mathbf{y}; \Lambda)$, appearing in the Finite Difference Monte Carlo risk (17) and gradient of the risk (19) estimators, are very cheap to compute, involving $\mathcal{O}(N)$ operations.

4.3.4 Evaluation of $\text{Tr}(\mathbf{A} \mathcal{S} \mathbf{A}^*)$

The evaluation of risk estimate $\hat{R}_{\nu,\varepsilon}(\ell; \Lambda | \mathcal{S})$, at Step (40) of generalized SURE and SUGAR Algorithm 2 requires the computation of the trace of an $N \times N$ matrix, with N possibly of order 10^6 , e.g. in image processing.

In the present application, combining the structure of covariance matrix $\mathcal{S}$ (81) and the sparse expression of $\mathbf{A}$ (84), provides a compact expression of the third term of generalized SURE (17) detailed in Proposition 5, which can be evaluated with very little computational effort.

Proposition 5 (Third term of Stein Unbiased Risk Estimate). *Consider texture’s leader coefficients (58), whose covariance matrix $\mathcal{S}$ evidences the sparse structure described in (81). Define the linear operator $\mathbf{A}$ from Formula (9), using operator Φ (66) and projector Π (71). Then, the third term of Stein estimator of the risk (17) reads*

$$\text{Tr}(\mathbf{A} \mathcal{S} \mathbf{A}^*) = \frac{N/2}{(F_0 F_2 - F_1^2)^2} \left(\sum_{j,j'} \left(F_1^2 \mathcal{C}_j^{j'} - 2F_0 F_1 j' \mathcal{C}_j^{j'} + F_0^2 j j' \mathcal{C}_j^{j'} \right) \right), \quad (85)$$

where the quantities $\{F_\alpha, \alpha = 0, 1, 2\}$ are defined in (68) and $\mathcal{C}_j^{j'}$ denotes the covariance between scales 2^j and $2^{j'}$, as defined in (82).

Proof. Proof is postponed to Appendix D. $\square$

5 Hyperparameter tuning performance assessment

The aim of this section is to assess quantitatively, by means of numerical simulations, the performance in the estimation of the optimal hyperparameters. To that end, Section 5.1 will detail the numerical simulation set-up and Section 5.2 will concentrate on several algorithmic issues. Section 5.3 will show on the prominent role of covariance matrix $\mathcal{S}$, evaluating the impact of partial vs. full covariance matrix in Section 5.3.2 and comparing true vs. estimated covariance matrix in Section 5.3.3. Section 5.4 will further assess quantitatively how well optimal hyperparameters are estimated in the absence of available ground truth, with respect to different quality metrics.

5.1 Numerical simulation set-up

5.1.1 Textures

For sake of simplicity, we consider the two-region case $M = 2$, with elliptic mask displayed in Figure 1a. Synthetic textures of resolution $N_1 \times N_2 = 256 \times 256$, characterized by two attributes configurations:

- Configuration “**D**”, “difficult”, one realization being displayed in Figure 1b

$$\begin{aligned} (\overline{H}_1, \overline{\Sigma}_1^2) &= (0.5, 0.6) & (\text{background}), \\ (\overline{H}_2, \overline{\Sigma}_2^2) &= (0.75, 0.7) & (\text{central ellipse}). \end{aligned}$$

- Configuration “**E**”, “easy”, one realization being displayed in Figure 1c

$$\begin{aligned} (\overline{H}_1, \overline{\Sigma}_1^2) &= (0.5, 0.6) & (\text{background}), \\ (\overline{H}_2, \overline{\Sigma}_2^2) &= (0.9, 1.1) & (\text{central ellipse}). \end{aligned}$$

are generated from a Matlab routine designed by ourselves (see [51]).

5.1.2 Multiscale analysis

A 2D undecimated wavelet transform of the textured image is computed at scale 2^j , with mother wavelet obtained as a tensor product of 1D least asymmetric Daubechies wavelets, with 3 vanishing moments, see [46] for more details.

5.1.3 Performance evaluation

Following [51, 52], for a given textured image X , and the derived log-leaders $(\ell_j)_{j=j_1}^{j_2}$, two performance indices are used:

- The *one-sample quadratic risk* on local regularity, computed from *one sample* of log-leaders ℓ computed on the single image X

$$\mathcal{R}(\ell; \Lambda) \triangleq \left\| \hat{\mathbf{h}}(\ell; \Lambda) - \bar{\mathbf{h}} \right\|_2^2, \quad (86)$$

with estimator $\hat{\mathbf{h}}(\ell; \Lambda)$ defined in (60) and ground truth $\bar{\mathbf{h}}$ defined in (57).

- The *segmentation error*, defined as the percentage of incorrectly classified pixels

$$\mathcal{P}(\ell; \Lambda) \triangleq \left| \bar{\Omega}_1 \cap \hat{\Omega}_2(\ell; \Lambda) \right| + \left| \hat{\Omega}_1(\ell; \Lambda) \cap \bar{\Omega}_2 \right|, \quad (87)$$

where $\cup_m \hat{\Omega}_m(\ell; \Lambda)$ is the estimated partition (62), obtain from TV-based texture segmentation, as described in Section 4.1.3.

Remark 10. By definition of the quadratic risk (70) and one-sample quadratic risk (86), $\mathbb{E}_{\mathcal{L}} \mathcal{R}(\ell; \Lambda) = R[\hat{\mathbf{h}}](\Lambda)$. In practice however, only one realization of ℓ is available, hence the quadratic risk $R[\hat{\mathbf{h}}](\Lambda)$ is not accessible. Thus, in the following experiments, the *one-sample quadratic risk* $\mathcal{R}(\ell; \Lambda)$, defined in (86), is used as a reference to which Stein risk estimator $\hat{R}_{\nu, \varepsilon}(\ell; \Lambda | \mathcal{S})$ will be compared.

5.2 Algorithmic set-up

5.2.1 Primal dual with iterative differentiation

Problem (60) is solved using the accelerated primal-dual algorithm 1, with primal variable $\mathbf{x} \triangleq (\mathbf{h}, \mathbf{v})$, taking advantage of strong-convexity of the data fidelity term. The maximal number of iterations is set to $K_{\max} = 5 \cdot 10^5$, and a threshold on the normalized duality gap is set to 10^{-4} (see [51]).

5.2.2 Scaling range

The estimation of piecewise constant local attributes requires to focus on fine scales. Thus, ideally, the least square term (59) would involve the two finest scales of the multiscale representation, and range from $j_1 = 1$ to $j_2 = 2$. Yet, the efficiency of acceleration strategy of Algorithm 1 increases with the strong-convexity modulus γ (73), displayed in Table 1, which is observed to increase with j_2 , as $j_1 = 1$ is fixed. Thus, a trade-off between locality and convergence speed leads to select $j_2 = 3$.

5.2.3 Finite Difference Monte Carlo parameters

The Monte Carlo vector $\varepsilon \in \mathbb{R}^P$, $P = JN_1N_2$, is drawn randomly, according to a i.i.d. normalized Gaussian $\mathcal{N}(\mathbf{0}_P, \mathbf{I}_P)$. We adapt the heuristic of [24] or the Finite Difference step ν to the case of correlated noise as

$$\nu = \frac{2}{P^\alpha} \max \left(\sqrt{C_j^j}, j \in \{1, \dots, J\} \right), \quad \alpha = 0.3, \quad (88)$$

	$j_2 = 2$	$j_2 = 3$	$j_2 = 4$	$j_2 = 5$	$j_2 = 6$
γ	0.29	0.72	1.20	1.69	2.20

Table 1: Strong-convexity modulus γ of data-fidelity term of (72), computed from Formula (73), for fixed $j_1 = 1$ and varied j_2 . The bold entry correspond to the range of scales used in the experiments of Sections 5.

where C_j^j is the variance of the log-leaders ℓ_j at scale 2^j .

The derivatives with respect to hyperparameters of the estimates, $\partial_{\mathbf{\Lambda}} \hat{\mathbf{h}}$, $\partial_{\mathbf{\Lambda}} \hat{\mathbf{v}}$, are obtained by iterative differentiation of primal dual algorithm, customized to texture segmentation in Appendix 1.

5.2.4 BFGS quasi-Newton initialization and parameters

To perform the risk minimization sketched in Algorithm 3, we used the GRadient-based Algorithm for Non-Smooth Optimization, implemented in GRANSO toolbox³, from the BFGS quasi-Newton algorithm proposed in [22]. It consists of a low memory BFGS algorithm with box constraints, enabling to enforce positive λ_h and λ_v . The maximal number of iterations of BFGS Algorithm 3 is set to $K_{\max} = 250$, while the stopping criterion on the gradient norm is set to 10^{-6} . As mentioned in Section 3.3, the initialization of quasi-Newton algorithms might drastically impact their convergence. Hence, we propose a model-based strategy for initializing $\mathbf{\Lambda}$ and $\mathbf{H}$. The initialization of λ_h and λ_v is performed by balancing the data fidelity term and the penalization appearing of functional (72). The data fidelity term grows like the variance of the noise

$$\mathbb{E} \sum_{j=j_1}^{j_2} \|j\bar{\mathbf{h}} + \bar{\mathbf{v}} - \ell_j\|_2^2 = \text{tr}(\mathbf{S}), \quad (89)$$

and the penalization term can be evaluated using $(\hat{\mathbf{h}}_{\text{LR}}, \hat{\mathbf{v}}_{\text{LR}})$ introduced in (59). Thus, the initial hyperparameters $\mathbf{\Lambda}$ for BFGS Algorithm 3 are set to

$$\mathbf{\Lambda}^{[0]} = (\lambda_h^{[0]}, \lambda_v^{[0]}), \quad \text{where} \quad \lambda_h^{[0]} = \frac{\text{tr}(\mathbf{S})}{2 \text{TV}(\hat{\mathbf{h}}_{\text{LR}}(\ell))}, \quad \text{and} \quad \lambda_v^{[0]} = \frac{\text{tr}(\mathbf{S})}{2 \text{TV}(\hat{\mathbf{v}}_{\text{LR}}(\ell))}. \quad (90)$$

The *inverse* Hessian matrix $\mathbf{H}^{[0]} \in \mathbb{R}^{2 \times 2}$, is initialized to enforce $\mathbf{\Lambda}^{[1]} = (1 \pm \kappa)\mathbf{\Lambda}^{[0]}$. It is chosen diagonal with coefficients

$$\mathbf{H}^{[0]} = \text{diag} \left(\left| \frac{\kappa \lambda_h^{[0]}}{\partial_{\lambda_h} \hat{R}_{\nu, \epsilon}(\ell; \mathbf{\Lambda}^{[0]} | \mathbf{S})} \right|, \left| \frac{\kappa \lambda_v^{[0]}}{\partial_{\lambda_v} \hat{R}_{\nu, \epsilon}(\ell; \mathbf{\Lambda}^{[0]} | \mathbf{S})} \right| \right). \quad (91)$$

In practice, we used $\kappa = 0.5$ for all experiments. It is observed that this choice of $\mathbf{H}^{[0]}$ avoids the first iteration falling away from natural hyperparameters scaling (90), which would induce huge computational cost to reach to optimal hyperparameters.

³<http://www.timitchell.com/software/GRANSO/>

5.3 Covariance of *leaders*

5.3.1 Covariance estimation procedure

No closed-form formula exists to compute exactly the covariance matrix $\mathbf{S}$ from the texture's attributes. Hence, from *one* sample ℓ , computed from a single texture X , the *estimated* covariance matrix, denoted $\hat{\mathbf{S}}$, is computed using classic sample covariance estimator:

$$\hat{\mathbf{S}}_{j,\underline{n}}^{j',\underline{n}'} \triangleq \frac{1}{|\Omega|} \sum_{\underline{n} \in \Omega} \ell_{j,\underline{n}} \ell_{j',\underline{n}+\delta\underline{n}} - \left(\frac{1}{|\Omega|} \sum_{\underline{n} \in \Omega} \ell_{j,\underline{n}} \right) \left(\frac{1}{|\Omega|} \sum_{\underline{n} \in \Omega} \ell_{j',\underline{n}} \right), \quad (92)$$

for spatial lag $\delta\underline{n} \triangleq \underline{n}' - \underline{n}$, leading to inter-scale covariance

$$\hat{\mathcal{C}}_j^{j'} = \hat{\mathbf{S}}_{j,\underline{n}}^{j',\underline{n}} \quad (93)$$

and spatial correlations

$$\hat{\Xi}_j^{j'}(\delta\underline{n}) = \frac{\hat{\mathbf{S}}_{j,\underline{n}}^{j',\underline{n}'}}{\hat{\mathcal{C}}_j^{j'}}. \quad (94)$$

Then, for Textures “D” and “E”, a *true* covariance matrix $\mathbf{S}$ is obtained numerically by averaging the above estimated covariance matrix $\hat{\mathbf{S}}^{(q)}$ over $Q = 5000$ texture samples as:

$$\mathbf{S} \triangleq \left\langle \hat{\mathbf{S}}^{(q)} \right\rangle_{q=1}^Q, \quad (95)$$

the samples being generated with the mathematical model of [51].

5.3.2 Impact of partial versus full covariance on estimated risk

We now assess the impact of using two partial versions of the *full true* covariance matrix $\mathbf{S}$, described in (95):

1. *Variance* matrix $\mathbf{S}_{\text{var}}$ neglecting both inter-scale and spatial correlations, reduces to the variances $\mathcal{C}_j^j$ of the ℓ_j 's, and hence is diagonal

$$S_{\text{var},j,\underline{n}}^{j',\underline{n}'} = \mathcal{C}_j^j \delta_{j,j'} \delta_{\underline{n},\underline{n}'}. \quad (96)$$

2. *Inter-scale* covariance matrix $\mathbf{S}_{\text{int}}$, neglecting spatial correlations, reduces to cross-correlations $\mathcal{C}_j^{j'}$ between the ℓ_j 's and the $\ell_{j'}$'s at same location

$$S_{\text{int},j,\underline{n}}^{j',\underline{n}'} = \mathcal{C}_j^{j'} \delta_{\underline{n},\underline{n}'}. \quad (97)$$

For texture “D”, both $\hat{R}_{\nu,\epsilon}(\ell; \mathbf{\Lambda} | \mathbf{S}_{\text{var}})$ (Fig. 3e) and $\hat{R}_{\nu,\epsilon}(\ell; \mathbf{\Lambda} | \mathbf{S}_{\text{int}})$ (Fig. 3i) fail to reproduce $\mathcal{R}(\ell; \mathbf{\Lambda})$ (Fig. 3a). Hence, the selected hyperparameters $\hat{\mathbf{\Lambda}}_{\nu,\epsilon}^\dagger(\ell | \mathbf{S}_{\text{var}})$ (‘□’) and $\hat{\mathbf{\Lambda}}_{\nu,\epsilon}^\dagger(\ell | \mathbf{S}_{\text{int}})$ (‘◇’) do not coincide with the optimal $\mathbf{\Lambda}_{\mathcal{R}}$ (‘+’). The corresponding segmentations, $T\hat{\mathbf{h}}(\ell; \hat{\mathbf{\Lambda}}_{\nu,\epsilon}^\dagger(\ell | \mathbf{S}_{\text{var}}))$ (Fig. 3f) and $T\hat{\mathbf{h}}(\ell; \hat{\mathbf{\Lambda}}_{\nu,\epsilon}^\dagger(\ell | \mathbf{S}_{\text{int}}))$

(Fig. 3j), differ significantly from the targeted $T\hat{\mathbf{h}}(\ell; \mathbf{\Lambda}_{\mathcal{R}})$ (Fig. 3b). On the opposite, $\hat{R}_{\nu,\epsilon}(\ell; \mathbf{\Lambda}|\mathcal{S})$ (Fig. 3m) perfectly matches $\mathcal{R}(\ell; \mathbf{\Lambda})$ (Fig. 3a). Thanks to the exact computation of the constant term $\text{Tr}(\mathbf{A}\mathcal{S}\mathbf{A}^*)$ in Proposition 5, the order of magnitude $\mathcal{R}(\ell; \mathbf{\Lambda})$ is well reproduced by $\hat{R}_{\nu,\epsilon}(\ell; \mathbf{\Lambda}|\mathcal{S})$, as observed on the colorbars in Figure 3. Further, $\hat{\mathbf{\Lambda}}_{\nu,\epsilon}^\dagger(\ell|\mathcal{S}_{\text{int}})$ ($\triangle$) coincides with $\mathbf{\Lambda}_{\mathcal{R}}$ ($+$), leading to segmentation $T\hat{\mathbf{h}}(\ell; \hat{\mathbf{\Lambda}}_{\nu,\epsilon}^\dagger(\ell|\mathcal{S}))$ (Fig. 3n) similar to $T\hat{\mathbf{h}}(\ell; \mathbf{\Lambda}_{\mathcal{R}})$ (Fig. 3b). Similar observations can be made for Texture “E” at columns 3, 4 of Figure 3.

Altogether, these two examples illustrate that the *full* covariance is necessary so that $\hat{R}_{\nu,\epsilon}(\ell; \mathbf{\Lambda}|\mathcal{S})$ provides an accurate estimate of $\mathcal{R}(\ell; \mathbf{\Lambda})$. Moreover, $\mathbf{\Lambda}_{\mathcal{R}}$ appears to be well approximated by the optimal hyperparameters $\hat{\mathbf{\Lambda}}_{\nu,\epsilon}^\dagger(\ell|\mathcal{S})$, obtained using *full* covariance.

5.3.3 Impact of estimating the covariance matrix

In practice, one has access to only the *estimated* covariance matrix $\hat{\mathcal{S}}$. This Section compares generalized SURE computed from *estimated* covariance $\hat{\mathcal{S}}$ to SURE computed assuming the knowledge of *true* covariance $\mathcal{S}$.

For Texture “D”, $\hat{R}_{\nu,\epsilon}(\ell; \mathbf{\Lambda}|\hat{\mathcal{S}})$ (Figure 4b) is identical to $\hat{R}_{\nu,\epsilon}(\ell; \mathbf{\Lambda}|\mathcal{S})$ (Figure 4a). Further, optimal hyperparameters $\hat{\mathbf{\Lambda}}_{\nu,\epsilon}^\dagger(\ell|\hat{\mathcal{S}})$ ($\triangle$) perfectly matches $\hat{\mathbf{\Lambda}}_{\nu,\epsilon}^\dagger(\ell|\mathcal{S})$ ($\triangle$) and lead to similar segmentations, $T\hat{\mathbf{h}}(\ell; \hat{\mathbf{\Lambda}}_{\nu,\epsilon}^\dagger(\ell|\hat{\mathcal{S}}))$ (Figure 4f) and $T\hat{\mathbf{h}}(\ell; \hat{\mathbf{\Lambda}}_{\nu,\epsilon}^\dagger(\ell|\mathcal{S}))$ (Figure 4e). These observations are precisely quantified in Table 2 in term of values of $\mathcal{R}(\ell; \mathbf{\Lambda})$ and percentage of misclassified pixels. The same observations can be made for Texture “E”.

Altogether, Figure 4 and the quantitative results reported in Table 2 show that $\hat{R}_{\nu,\epsilon}(\ell; \mathbf{\Lambda}|\hat{\mathcal{S}})$ provides an accurate estimate of $\mathcal{R}(\ell; \mathbf{\Lambda})$, and that $\hat{\mathbf{\Lambda}}_{\nu,\epsilon}^\dagger(\ell|\hat{\mathcal{S}})$ is a good estimate of $\mathbf{\Lambda}_{\mathcal{R}}$.

	Texture “D”		Texture “E”	
Hyperparameter $\mathbf{\Lambda}$	$\mathcal{R}(\ell; \mathbf{\Lambda})$	$\mathcal{P}(\ell; \mathbf{\Lambda})$	$\mathcal{R}(\ell; \mathbf{\Lambda})$	$\mathcal{P}(\ell; \mathbf{\Lambda})$
$\mathbf{\Lambda}_{\mathcal{R}}$ ‘+’	$2.32 \cdot 10^3$	7.79%	$2.66 \cdot 10^3$	5.34%
$\hat{\mathbf{\Lambda}}_{\nu,\epsilon}^\dagger(\ell \mathcal{S})$ ‘ $\triangle$ ’	$2.35 \cdot 10^3$	5.51%	$2.83 \cdot 10^3$	9.58%
$\hat{\mathbf{\Lambda}}_{\nu,\epsilon}^\dagger(\ell \hat{\mathcal{S}})$ ‘ $\triangle$ ’	$2.35 \cdot 10^3$	5.51%	$2.96 \cdot 10^3$	4.61%
$\hat{\mathbf{\Lambda}}_{\nu,\epsilon}^{\text{BFGS}}(\ell \mathcal{S})$ ‘ ∇ ’	$2.36 \cdot 10^3$	4.66%	$2.83 \cdot 10^3$	3.71%
$\hat{\mathbf{\Lambda}}_{\nu,\epsilon}^{\text{BFGS}}(\ell \hat{\mathcal{S}})$ ‘ ∇ ’	$2.36 \cdot 10^3$	6.22%	$2.83 \cdot 10^3$	3.27%

Table 2: Grid search v.s. BFGS Algorithm 3 performance in term of quadratic error $\mathcal{R}(\ell; \mathbf{\Lambda})$ and segmentation error $\mathcal{P}(\ell; \mathbf{\Lambda})$ for the two different Textures “D” and “E”.

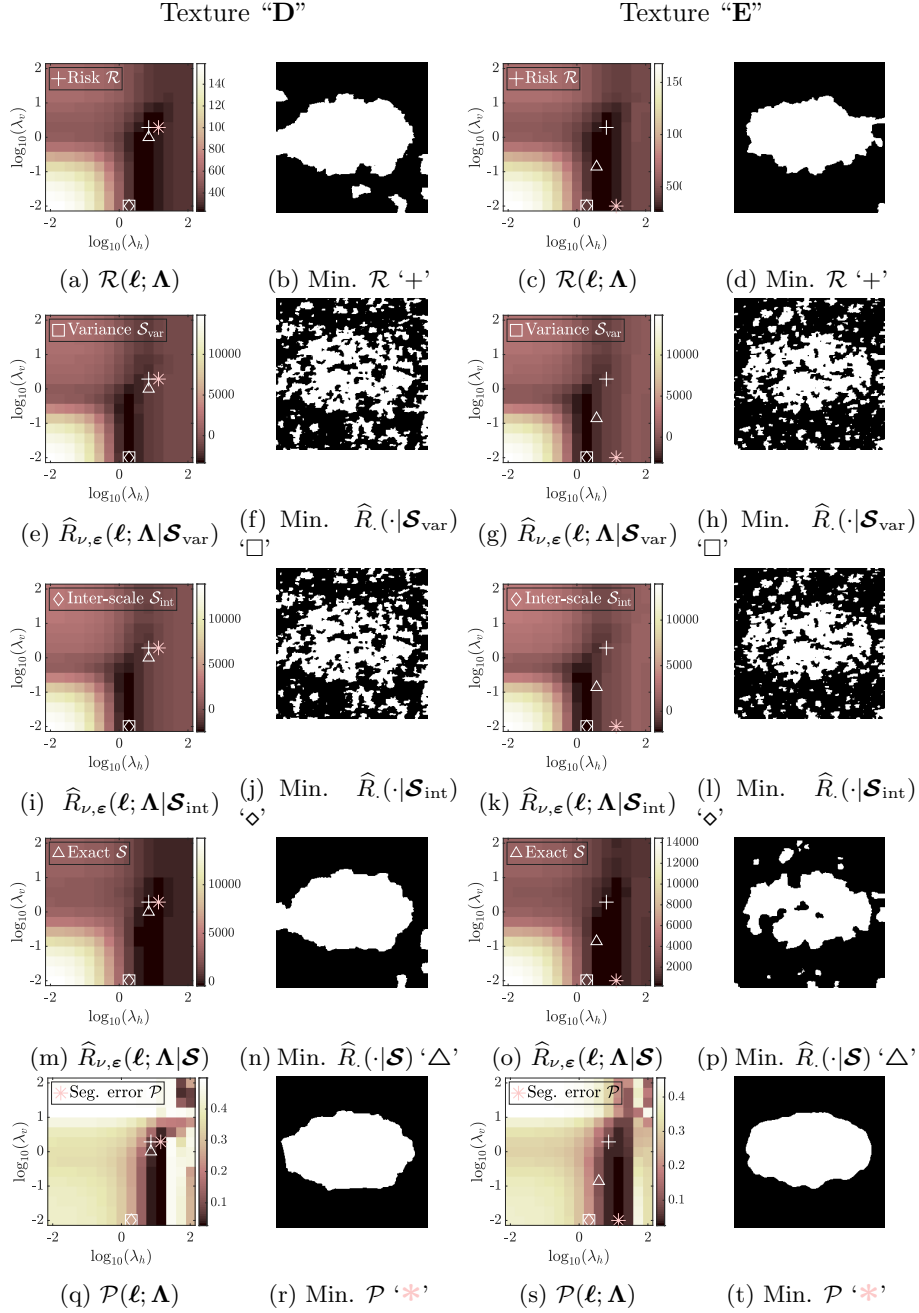


Figure 3: Error maps for TV-based texture segmentation on a grid of $\Lambda = (\lambda_h, \lambda_v)$, and segmentation obtained with associated optimal hyperparameters for piecewise Textures “D” (column 1, 2) and “E” (column 3, 4). Estimated risks $\hat{R}_{\nu, \epsilon}(\ell; \Lambda | \mathcal{S})$ computed either with variance matrix $\mathcal{S}_{\text{var}}$ (second row), inter-scale covariance matrix $\mathcal{S}_{\text{int}}$ (third row), or full covariance matrix $\mathcal{S}$ (fourth row) are compared.

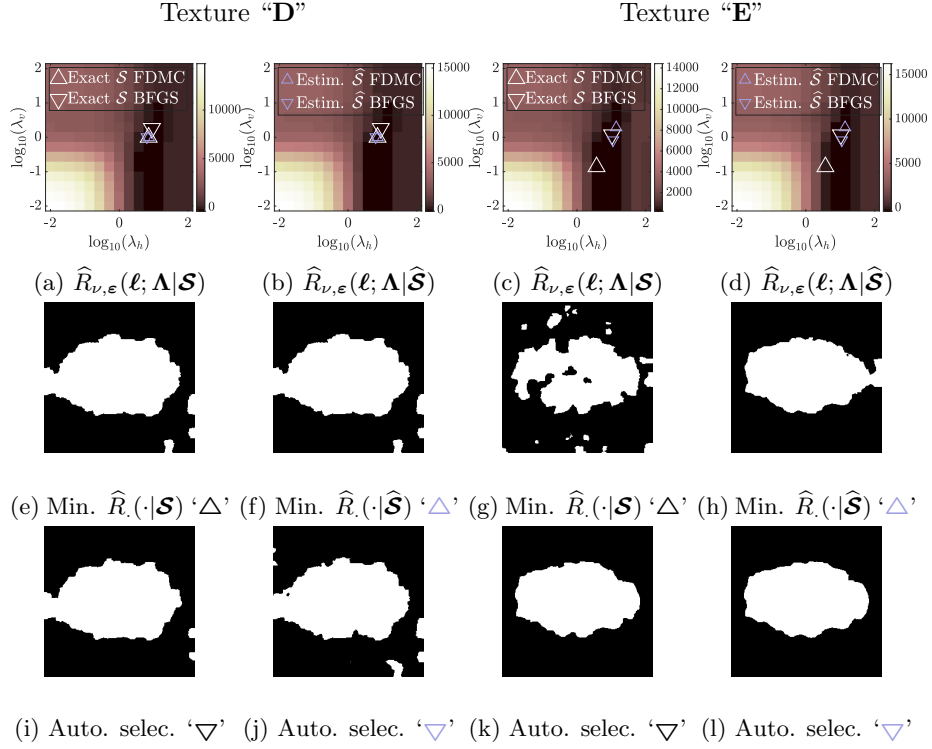


Figure 4: Generalized SURE computed either from *true* covariance matrix $\mathcal{S}$ (95), or from *estimated* covariance matrix (92) for Textures “D” and “E” (first row). Segmentations obtained minimizing the above generalized SURE (second row). Segmentations obtained with automated selection of hyperparameters from Algorithm 3, using generalized SUGAR with either *true* covariance matrix or *estimated* covariance matrix (third row).

5.4 Automated selection of hyperparameters

Section 5.3 has shown the relevance of Algorithm 3 by comparing its performance against those obtained from a grid search on hyperparameters Λ . Section 5.4 will now test the practical effectiveness of the proposed procedure by assessing the convergence of the quasi-Newton algorithm and corresponding performance in hyperparameter selection and segmentation, avoiding the recourse to any ground truth and hence to the greedy and unfeasible grid search.

5.4.1 Effective convergence of quasi-Newton Algorithm

The convergence of quasi-Newton Algorithm 3 is assessed empirically comparing automatically selected hyperparameters $\hat{\Lambda}_{\nu, \epsilon}^{\text{BFGS}}$ with optimal hyperparameters found from exhaustive grid search $\hat{\Lambda}_{\nu, \epsilon}^{\dagger}$.

Figures 4a and 4b illustrate that $\hat{\Lambda}_{\nu, \epsilon}^{\text{BFGS}}(\ell|\mathcal{S})$ (‘ ∇ ’) and $\hat{\Lambda}_{\nu, \epsilon}^{\text{BFGS}}(\ell|\hat{\mathcal{S}})$ (‘ ∇ ’) respectively match $\hat{\Lambda}_{\nu, \epsilon}^{\dagger}(\ell|\mathcal{S})$ (‘ Δ ’) and $\hat{\Lambda}_{\nu, \epsilon}^{\dagger}(\ell|\hat{\mathcal{S}})$ (‘ Δ ’) in the case of Tex-

ture “D”. Similar conclusions can be drawn from Figures 4c and 4d for Texture “E”.

Figure 4 and the quantitative results provided in Table 2 show the convergence of Algorithm 3 using $\mathcal{S}$ (resp. $\hat{\mathcal{S}}$) toward the minimum of $\hat{R}_{\nu,\epsilon}(\ell; \Lambda|\mathcal{S})$ (resp. $\hat{R}_{\nu,\epsilon}(\ell; \Lambda|\hat{\mathcal{S}})$).

In term of computational cost, Algorithm 3 requires an average of 40 calls of Algorithm 1, compared to 225 calls needed to perform grid search at Section 5.3.2.

5.4.2 Automated selection of Λ and segmentation performance

Ten realizations of Textures “D” and “E” are generated following the procedure described in Section 5.1.1. For each of them, Algorithm 3 is run twice, first using $\mathcal{S}$ and second using $\hat{\mathcal{S}}$.

Since here no grid search is performed, the minimum value of quadratic risk is unknown. The performance will hence be measured in terms of *normalized one-sample quadratic risk* $\tilde{\mathcal{R}}$ defined as

$$\tilde{\mathcal{R}}(\ell|\mathcal{S}) = \frac{\mathcal{R}(\ell; \hat{\Lambda}_{\nu,\epsilon}^{\text{BFGS}}(\ell|\mathcal{S}))}{\|\hat{\mathbf{h}}_{\text{LR}}(\ell) - \bar{\mathbf{h}}\|_2^2} = \frac{\|\hat{\mathbf{h}}_{\nu,\epsilon}^{\text{BFGS}}(\ell|\mathcal{S}) - \bar{\mathbf{h}}\|_2^2}{\|\hat{\mathbf{h}}_{\text{LR}}(\ell) - \bar{\mathbf{h}}\|_2^2}, \quad (98)$$

measuring the improvement of the estimation achieved using TV-based texture segmentation (60) with hyperparameters automatically selected by Algorithm 3, compared to the classical least square estimate $\hat{\mathbf{h}}_{\text{LR}}$.

Averaged performance over ten realizations, presented in Table 3, show that the quadratic risk $\mathcal{R}$ obtained is decrease by a factor of 16 for Texture “D” and of 14 for Texture “E”. The corresponding *segmentation error* is as low as 6% for Texture “D”, and 3% for Texture “E”. Further, the use of *estimated* covariance matrix does not degrade achieved performance compare to using *true* covariance matrix.

Hence, Algorithm 3, using the *estimated* covariance $\hat{\mathcal{S}}$, computed from (92), provides an efficient, parameter-free, automated and data-driven texture segmentation procedure.

	Texture “D”		Texture “E”	
Covariance matrix	$\mathcal{S}$	$\hat{\mathcal{S}}$	$\mathcal{S}$	$\hat{\mathcal{S}}$
$\tilde{\mathcal{R}}(\ell \cdot)$	0.060 ± 0.003	0.057 ± 0.002	0.071 ± 0.003	0.073 ± 0.004
$\mathcal{P}(\ell; \hat{\Lambda}_{\nu,\epsilon}^{\text{BFGS}}(\ell \cdot))$ (%)	5.4 ± 0.7	6.8 ± 1.5	3.3 ± 0.7	2.8 ± 0.3

Table 3: Averaged performance of TV-based texture segmentation with automated selection of hyperparameters.

6 Conclusion

This work was focused on devising a procedure for the automated selection of the hyperparameters of parametric estimators, such as e.g., parametric linear filtering or penalized least squares. The main result obtained here consists of a theoretically grounded and practical operational fully-automated data driven procedure, that requires neither ground truth nor expert-based knowledge and work satisfactorily even when applied to a single observation of data.

To that end, Stein Unbiased Risk Estimator (SURE) was rewritten to account for additive correlated Gaussian noise, with any covariance structure. The main contribution compared to state-of-the-art procedure relies on including the covariance matrix of the noise only in SURE, rather than in the data fidelity term. The benefit is twofold: handling with a strongly convex function when Penalized Least Square is considered, and avoiding costly, if not intractable, inversion of the covariance matrix. Differentiating this Generalized SURE with respect to hyperparameters, an estimator for the risk gradient was designed, permitting to propose a Generalized Finite Difference Monte Carlo Stein Unbiased GrAdient Risk (SUGAR) estimate. The asymptotic unbiasedness of Generalized SUGAR was assessed theoretically, based on regularity assumptions on the parametric estimator.

Further, the case of sequential parametric estimators is discussed in depth in the case of primal-dual minimization scheme for Penalized Least Squares and a differentiated scheme is derived.

Embedding Generalized SURE and SUGAR into a quasi-Newton algorithm enabled to perform an automated risk minimization. An explicit algorithm permitting to implement the minimization was proposed.

To assess the performance of this automated hyperparameter selection procedure devised in a general setting, it has been customized to the specific problem of texture segmentation, based on multiscale descriptors (wavelet leaders) and nonsmooth Total-Variation based penalization. This problem is uneasy because observations are in nature multiscale, with inhomogeneous variance across scales and correlations both across scales and in space at each scale. Further, variances and correlations are unknown and need to be estimated directly from data.

Numerical simulations, conducted on ten realizations of synthetic piecewise fractal textures, permitted to show that the proposed strategy yield satisfactory performance in selecting automatically the penalization hyperparameter, leading to excellent texture segmentation, with no ad-hoc (or expert-based) tuning and without prior knowledge for ground truth, and using one-sample estimate of the covariance matrix.

The corresponding MATLAB routines, developed by ourselves and implementing these tools, ready for applications to real-world texture segmentation, where hyperparameter tuning constitutes an on-going hot topic, will be made publicly available to the research community in a documented toolbox at the time of publication.

A Proof of Theorem 1

Proof. For ease of computation we first define the *predictor* in Definition 4 and the ground truth *prediction* in Definition 5.

Definition 4 (*Predictor*). From the estimator of underlying features $\hat{\mathbf{x}}(\mathbf{y}; \Lambda)$ one can equivalently consider a *prediction* estimator

$$\hat{\mathbf{y}}(\mathbf{y}; \Lambda) \triangleq \Phi \hat{\mathbf{x}}(\mathbf{y}; \Lambda). \quad (99)$$

Indeed, from Assumption 2, $\Phi^* \Phi$ is invertible, and the relation (99) can be inverted computing

$$\hat{\mathbf{x}}(\mathbf{y}; \Lambda) = (\Phi^* \Phi)^{-1} \Phi^* \hat{\mathbf{y}}(\mathbf{y}; \Lambda). \quad (100)$$

Definition 5 (*Prediction* ground truth). The noise-free observation writes

$$\bar{\mathbf{y}} \triangleq \mathbb{E}_\zeta \mathbf{y} = \Phi \bar{\mathbf{x}}. \quad (101)$$

Thus, the quadratic risk defined in (10) can be expressed using operator $\mathbf{A}$ defined in (9) as

$$\begin{aligned} R[\hat{\mathbf{x}}](\Lambda) &= \mathbb{E}_\zeta \|\Pi \hat{\mathbf{x}}(\mathbf{y}; \Lambda) - \Pi \mathbf{x}\|_2^2 = \mathbb{E}_\zeta \|\Pi (\Phi^* \Phi)^{-1} \Phi^* (\hat{\mathbf{y}}(\mathbf{y}; \Lambda) - \bar{\mathbf{y}})\|_2^2, \\ &\stackrel{(9)}{=} \mathbb{E}_\zeta \|\mathbf{A} (\hat{\mathbf{y}}(\mathbf{y}; \Lambda) - \bar{\mathbf{y}})\|_2^2 \end{aligned} \quad (102)$$

which will be easier to manipulate in the following when expressed in term of noise-free (or noisy) observations $\bar{\mathbf{y}}$ (or $\mathbf{y}$) and *prediction* $\hat{\mathbf{y}}$.

By construction, the matrix $\mathbf{A}$, defined in (9), performs both:

- The projection on the interest subspace $\mathcal{I}$ of $\mathcal{H}$ via the linear operator Π .
- The transition from predicted quantities $\hat{\mathbf{y}}$ to estimated features $\hat{\mathbf{x}}$, making use of relation (100).

From now, for sake of simplicity, we make implicit the dependency of $\hat{\mathbf{x}}$ in $(\mathbf{y}; \Lambda)$. From the model (1) and the Assumption 1 on the noise probability distribution, one directly derive two useful relations:

$$\mathbb{E}_\zeta \|\mathbf{A} (\mathbf{y} - \bar{\mathbf{y}})\|_2^2 \stackrel{(101)}{=} \mathbb{E}_\zeta \|\mathbf{A} (\mathbf{y} - \Phi \bar{\mathbf{x}})\|_2^2 = \mathbb{E}_\zeta \|\mathbf{A} \zeta\|_2^2 \stackrel{\text{Hyp. 1}}{=} \text{Tr}(\mathbf{A} \mathbf{S} \mathbf{A}^*), \quad (103)$$

$$\mathbb{E}_\zeta \langle \mathbf{A} \mathbf{y}, \mathbf{A} (\mathbf{y} - \bar{\mathbf{y}}) \rangle \stackrel{\mathbb{E}(\mathbf{y} - \bar{\mathbf{y}}) = 0}{=} \mathbb{E}_\zeta \langle \mathbf{A} (\mathbf{y} - \bar{\mathbf{y}}), \mathbf{A} (\mathbf{y} - \bar{\mathbf{y}}) \rangle \stackrel{\text{Hyp. 1}}{=} \text{Tr}(\mathbf{A} \mathbf{S} \mathbf{A}^*). \quad (104)$$

Thus the risk can be expanded as

$$\begin{aligned}
R[\hat{\mathbf{x}}](\mathbf{\Lambda}) &\triangleq \mathbb{E}_{\boldsymbol{\zeta}} \|\mathbf{A}(\hat{\mathbf{y}} - \bar{\mathbf{y}})\|_2^2 \\
&= \mathbb{E}_{\boldsymbol{\zeta}} \left[\|\mathbf{A}(\hat{\mathbf{y}} - \mathbf{y})\|_2^2 + \|\mathbf{A}(\mathbf{y} - \bar{\mathbf{y}})\|_2^2 + 2\langle \mathbf{A}(\hat{\mathbf{y}} - \mathbf{y}), \mathbf{A}(\mathbf{y} - \bar{\mathbf{y}}) \rangle \right] \\
&\stackrel{(103)}{=} \mathbb{E}_{\boldsymbol{\zeta}} \left[\|\mathbf{A}(\hat{\mathbf{y}} - \mathbf{y})\|_2^2 + 2\langle \mathbf{A}\hat{\mathbf{y}}, \mathbf{A}(\mathbf{y} - \bar{\mathbf{y}}) \rangle - 2\langle \mathbf{A}\mathbf{y}, \mathbf{A}(\mathbf{y} - \bar{\mathbf{y}}) \rangle \right] + \text{Tr}(\mathbf{A}\mathbf{S}\mathbf{A}^*) \\
&\stackrel{(104)}{=} \mathbb{E}_{\boldsymbol{\zeta}} \left[\|\mathbf{A}(\hat{\mathbf{y}} - \mathbf{y})\|_2^2 + 2\langle \mathbf{A}\hat{\mathbf{y}}, \mathbf{A}(\mathbf{y} - \bar{\mathbf{y}}) \rangle \right] - 2\text{Tr}(\mathbf{A}\mathbf{S}\mathbf{A}^*) + \text{Tr}(\mathbf{A}\mathbf{S}\mathbf{A}^*) \\
&= \mathbb{E}_{\boldsymbol{\zeta}} \left[\|\mathbf{A}(\hat{\mathbf{y}} - \mathbf{y})\|_2^2 + 2\langle \mathbf{A}\hat{\mathbf{y}}, \mathbf{A}(\mathbf{y} - \bar{\mathbf{y}}) \rangle \right] - \text{Tr}(\mathbf{A}\mathbf{S}\mathbf{A}^*) \\
&= \mathbb{E}_{\boldsymbol{\zeta}} \left[\|\mathbf{A}(\Phi\hat{\mathbf{x}} - \mathbf{y})\|_2^2 + 2\langle \mathbf{A}^*\mathbf{A}\Phi\hat{\mathbf{x}}, (\mathbf{y} - \Phi\bar{\mathbf{x}}) \rangle \right] - \text{Tr}(\mathbf{A}\mathbf{S}\mathbf{A}^*) \\
&= \mathbb{E}_{\boldsymbol{\zeta}} \|\mathbf{A}(\Phi\hat{\mathbf{x}} - \mathbf{y})\|_2^2 + 2\mathbb{E}_{\boldsymbol{\zeta}} \langle \mathbf{A}^*\mathbf{A}\Phi\hat{\mathbf{x}}, \boldsymbol{\zeta} \rangle - \text{Tr}(\mathbf{A}\mathbf{S}\mathbf{A}^*) \\
&= \mathbb{E}_{\boldsymbol{\zeta}} \|\mathbf{A}(\Phi\hat{\mathbf{x}} - \mathbf{y})\|_2^2 + 2\mathbb{E}_{\boldsymbol{\zeta}} \langle \mathbf{A}^*\mathbf{\Pi}\hat{\mathbf{x}}, \boldsymbol{\zeta} \rangle - \text{Tr}(\mathbf{A}\mathbf{S}\mathbf{A}^*),
\end{aligned}$$

re-injecting the definition of $\mathbf{A}$ (9) in terms of Φ and $\mathbf{\Pi}$.

The second term, $\mathbb{E}_{\boldsymbol{\zeta}} \langle \mathbf{A}^*\mathbf{\Pi}\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}), \boldsymbol{\zeta} \rangle$, is called the *degrees of freedom* [29]. From Assumption 3 it is well-defined and writes

$$\mathbb{E}_{\boldsymbol{\zeta}} \langle \mathbf{A}^*\mathbf{\Pi}\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}), \boldsymbol{\zeta} \rangle = \quad (105)$$

$$\frac{1}{\sqrt{(2\pi)^P |\det(\mathbf{S})|}} \int \langle \mathbf{A}^*\mathbf{\Pi}\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}), \boldsymbol{\zeta} \rangle \exp\left(-\frac{\boldsymbol{\zeta}^* \mathbf{S}^{-1} \boldsymbol{\zeta}}{2}\right) d\boldsymbol{\zeta}, \quad (106)$$

hence requiring generalized Stein's lemma to be estimated⁴.

Because of the off-diagonal terms in $\mathbf{S}^{-1}$, the Integration by Parts (IP) required to transform (105) cannot be directly justified, thus Stein's lemma generalization to $\mathcal{G}$ -valued random variable $\boldsymbol{\zeta}$ is not straightforward. Hence we propose to first diagonalize $\mathbf{S}^{-1}$ (which is a symmetric matrix) in a orthonormal basis, obtaining

$$\mathbf{S}^{-1} = \mathbf{V}^* \mathbf{D} \mathbf{V},$$

with $\mathbf{V}$ an orthonormal matrix (which columns are eigenvectors of $\mathbf{S}^{-1}$) and $\mathbf{D} = \text{diag}(\beta_1, \dots, \beta_P)$ containing (positive) eigenvalues of $\mathbf{S}^{-1}$. Then, setting $\boldsymbol{\vartheta} = \mathbf{V}\boldsymbol{\zeta}$

$$\begin{aligned}
\mathbb{E}_{\boldsymbol{\vartheta}} \langle \mathbf{A}^*\mathbf{\Pi}\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}), \boldsymbol{\vartheta} \rangle &= \\
&\frac{1}{\sqrt{(2\pi)^P |\det(\mathbf{S})|}} \int \langle \mathbf{A}^*\mathbf{\Pi}\hat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}), \mathbf{V}^{-1}\boldsymbol{\vartheta} \rangle \exp\left(-\frac{\boldsymbol{\vartheta}^* \mathbf{D} \boldsymbol{\vartheta}}{2}\right) |\det(\mathbf{V}^{-1})| d\boldsymbol{\vartheta}.
\end{aligned}$$

with $\boldsymbol{\vartheta}^* \mathbf{D} \boldsymbol{\vartheta} = \sum_{p=1}^P \beta_p |\vartheta_p|^2$.

⁴ Stein's lemma states that, for a real random variable $\zeta \sim \mathcal{N}(0, \sigma^2)$, if $f : \mathbb{R} \rightarrow \mathbb{R}$ is a function such that both $\mathbb{E}_{\zeta}[\zeta f(\zeta)]$ and $\mathbb{E}_{\zeta}[f'(\zeta)]$ exist, then $\mathbb{E}_{\zeta}[\zeta f(\zeta)] = \sigma^2 \mathbb{E}_{\zeta}[f'(\zeta)]$. Its demonstration relies on appropriate integration by parts.

Since $\mathbf{V}$ is orthonormal: $\mathbf{V}^{-1} = \mathbf{V}^*$ and $|\det(\mathbf{V}^{-1})| = 1$, leading to

$$\begin{aligned}
& \mathbb{E}_{\boldsymbol{\vartheta}} \langle \mathbf{A}^* \Pi \hat{\mathbf{x}}(\mathbf{y}; \Lambda), \boldsymbol{\vartheta} \rangle \\
&= \frac{1}{\sqrt{(2\pi)^P |\det(\mathcal{S})|}} \int \langle \mathbf{V} \mathbf{A}^* \Pi \hat{\mathbf{x}}(\mathbf{y}; \Lambda), \boldsymbol{\vartheta} \rangle \exp\left(-\frac{\boldsymbol{\vartheta}^* \mathcal{D} \boldsymbol{\vartheta}}{2}\right) d\boldsymbol{\vartheta} \\
&= \frac{1}{\sqrt{(2\pi)^P |\det(\mathcal{S})|}} \int \sum_{p=1}^P (\mathbf{V} \mathbf{A}^* \Pi \hat{\mathbf{x}}(\mathbf{y}; \Lambda))_p \vartheta_p \exp\left(-\frac{\sum_{p=1}^P \beta_p |\vartheta_p|^2}{2}\right) d\vartheta_1 \dots d\vartheta_P \\
&\stackrel{(\text{IP})}{=} \mathbb{E}_{\boldsymbol{\vartheta}} \left[\sum_{p=1}^P \frac{1}{\beta_p} \frac{\partial (\mathbf{V} \mathbf{A}^* \Pi \hat{\mathbf{x}}(\mathbf{y}; \Lambda))_p}{\partial \vartheta_p} \right] \\
&= \mathbb{E}_{\boldsymbol{\vartheta}} \left[\text{Tr} \left(\mathcal{D}^{-1} \frac{\partial (\mathbf{V} \mathbf{A}^* \Pi \hat{\mathbf{x}}(\mathbf{y}; \Lambda))}{\partial \boldsymbol{\vartheta}} \right) \right],
\end{aligned}$$

where $\frac{\partial (\mathbf{V} \mathbf{A}^* \Pi \hat{\mathbf{x}}(\mathbf{y}; \Lambda))}{\partial \boldsymbol{\vartheta}}$ denotes the Jacobian matrix of $\mathbf{V} \mathbf{A}^* \Pi \hat{\mathbf{x}}(\mathbf{y}; \Lambda)$ with respect to the variable $\boldsymbol{\vartheta} \triangleq \mathbf{V}^{-1} \boldsymbol{\zeta}$. In order to go back to variable $\boldsymbol{\zeta}$, we make use of (1) relating $\mathbf{y}$ and $\boldsymbol{\zeta}$, and apply the reverse change of variable $\boldsymbol{\zeta} \triangleq \mathbf{V} \boldsymbol{\vartheta}$ and obtain

$$\frac{\partial (\mathbf{V} \mathbf{A}^* \Pi \hat{\mathbf{x}}(\mathbf{y}; \Lambda))}{\partial \boldsymbol{\vartheta}} = \mathbf{V} \frac{\partial (\mathbf{A}^* \Pi \hat{\mathbf{x}}(\mathbf{y}; \Lambda))}{\partial \boldsymbol{\zeta}} \mathbf{V}^{-1} = \mathbf{V} \frac{\partial (\mathbf{A}^* \Pi \hat{\mathbf{x}}(\mathbf{y}; \Lambda))}{\partial \mathbf{y}} \mathbf{V}^{-1},$$

because $\partial_{\boldsymbol{\zeta}} \mathbf{y} = \mathbf{I}_P$ (the identity matrix of size $P \times P$).

Using the cyclicity of trace and the fact that $\mathbf{V}$ is orthonormal, we finally obtain a closed-form expression of the degrees of freedom:

$$\begin{aligned}
\mathbb{E}_{\boldsymbol{\vartheta}} \langle \mathbf{A}^* \Pi \hat{\mathbf{x}}(\mathbf{y}; \Lambda), \boldsymbol{\vartheta} \rangle &= \mathbb{E}_{\boldsymbol{\zeta}} \left[\text{Tr} \left(\mathcal{D}^{-1} \mathbf{V} \frac{\partial (\mathbf{A}^* \Pi \hat{\mathbf{x}}(\mathbf{y}; \Lambda))}{\partial \mathbf{y}} \mathbf{V}^{-1} \right) \right], \\
&= \mathbb{E}_{\boldsymbol{\zeta}} \left[\text{Tr} \left(\mathbf{V}^{-1} \mathcal{D}^{-1} \mathbf{V} \frac{\partial (\mathbf{A}^* \Pi \hat{\mathbf{x}}(\mathbf{y}; \Lambda))}{\partial \mathbf{y}} \right) \right], \\
&= \mathbb{E}_{\boldsymbol{\zeta}} \left[\text{Tr} \left(\mathcal{S} \frac{\partial (\mathbf{A}^* \Pi \hat{\mathbf{x}}(\mathbf{y}; \Lambda))}{\partial \mathbf{y}} \right) \right], \\
&= \mathbb{E}_{\boldsymbol{\zeta}} [\text{Tr} (\mathcal{S} \mathbf{A}^* \Pi \partial_{\mathbf{y}} \hat{\mathbf{x}}(\mathbf{y}; \Lambda))]
\end{aligned}$$

□

B Finite Difference Monte Carlo SURE

Proof. First, remark that since $\mathbf{y} \mapsto \hat{\mathbf{x}}(\mathbf{y}; \Lambda)$ is Lipschitz continuous from Assumption 4, it is Lebesgue differentiable almost everywhere and its Lebesgue derivative equals its weak derivative almost everywhere. Then, based on Theorem 1, the only difficulty relies in dominating the degrees of freedom, since it is the only term depending on the Finite Difference step ν .

Applying successively both Monte Carlo and Finite Difference strategies presented in Section 2.4 we obtain

$$\begin{aligned}
\text{Tr} (\mathcal{S} \mathbf{A}^* \Pi \partial_{\mathbf{y}} \hat{\mathbf{x}}(\mathbf{y}; \Lambda)) &\stackrel{\text{Monte Carlo}}{=} \mathbb{E}_{\boldsymbol{\varepsilon}} \left\langle \mathcal{S} \mathbf{A}^* \Pi \frac{\partial \hat{\mathbf{x}}(\mathbf{y}; \Lambda)}{\partial \mathbf{y}} [\boldsymbol{\varepsilon}], \boldsymbol{\varepsilon} \right\rangle \\
&\stackrel{\text{Finite Difference}}{=} \mathbb{E}_{\boldsymbol{\varepsilon}} \left\langle \mathbf{A}^* \Pi \lim_{\nu \rightarrow 0} \frac{\hat{\mathbf{x}}(\mathbf{y} + \nu \boldsymbol{\varepsilon}; \Lambda) - \hat{\mathbf{x}}(\mathbf{y}; \Lambda)}{\nu}, \mathcal{S} \boldsymbol{\varepsilon} \right\rangle.
\end{aligned} \tag{107}$$

Making use of the centered normalized Gaussian probability density function of ε , the above expectation writes

$$\mathbb{E}_\varepsilon \left\langle \mathbf{A}^* \Pi \lim_{\nu \rightarrow 0} \frac{\widehat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \Lambda) - \widehat{\mathbf{x}}(\mathbf{y}; \Lambda)}{\nu}, \mathbf{S} \varepsilon \right\rangle \quad (108)$$

$$(\text{p.d.f. of } \varepsilon) = \int_{\mathbb{R}^P} \lim_{\nu \rightarrow 0} \left\langle \mathbf{A}^* \Pi \frac{\widehat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \Lambda) - \widehat{\mathbf{x}}(\mathbf{y}; \Lambda)}{\nu}, \mathbf{S} \varepsilon \right\rangle \frac{e^{-\frac{\|\varepsilon\|^2}{2}} d\varepsilon}{(2\pi)^{P/2}}$$

Then the following majorations hold

$$\left| \left\langle \mathbf{A}^* \Pi \frac{\widehat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \Lambda) - \widehat{\mathbf{x}}(\mathbf{y}; \Lambda)}{\nu}, \mathbf{S} \varepsilon \right\rangle \right| e^{-\frac{\|\varepsilon\|^2}{2}} \quad (109)$$

$$(\text{Cauchy-Schwarz}) \leq \left\| \mathbf{A}^* \Pi \frac{\widehat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \Lambda) - \widehat{\mathbf{x}}(\mathbf{y}; \Lambda)}{\nu} \right\| \|\mathbf{S} \varepsilon\| e^{-\frac{\|\varepsilon\|^2}{2}}$$

$$(\text{Bounded operators}) \leq \|\mathbf{A}^*\| \|\Pi\| \left\| \frac{\widehat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \Lambda) - \widehat{\mathbf{x}}(\mathbf{y}; \Lambda)}{\nu} \right\| \|\mathbf{S}\| \|\varepsilon\| e^{-\frac{\|\varepsilon\|^2}{2}}$$

$$(\text{Hyp. 4: } L_1\text{-Lipschitz}) \leq \|\mathbf{A}^*\| \|\Pi\| L_1 \|\varepsilon\| \|\mathbf{S}\| \|\varepsilon\| e^{-\frac{\|\varepsilon\|^2}{2}},$$

with $\|\varepsilon\|^2 e^{-\frac{\|\varepsilon\|^2}{2}}$ integrable over $\mathbb{R}^P$. Further, the domination being independent of ν the limit can be interchanged with the integral on variable ε which gives

$$\int_{\mathbb{R}^P} \lim_{\nu \rightarrow 0} \left\langle \mathbf{A}^* \Pi \frac{\widehat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \Lambda) - \widehat{\mathbf{x}}(\mathbf{y}; \Lambda)}{\nu}, \mathbf{S} \varepsilon \right\rangle \frac{e^{-\frac{\|\varepsilon\|^2}{2}} d\varepsilon}{(2\pi)^{P/2}} \quad (110)$$

$$= \lim_{\nu \rightarrow 0} \int_{\mathbb{R}^P} \left\langle \mathbf{A}^* \Pi \frac{\widehat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \Lambda) - \widehat{\mathbf{x}}(\mathbf{y}; \Lambda)}{\nu}, \mathbf{S} \varepsilon \right\rangle \frac{e^{-\frac{\|\varepsilon\|^2}{2}} d\varepsilon}{(2\pi)^{P/2}},$$

and

$$\begin{aligned} & \left| \lim_{\nu \rightarrow 0} \int_{\mathbb{R}^P} \left\langle \mathbf{A}^* \Pi \frac{\widehat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \Lambda) - \widehat{\mathbf{x}}(\mathbf{y}; \Lambda)}{\nu}, \mathbf{S} \varepsilon \right\rangle \frac{e^{-\frac{\|\varepsilon\|^2}{2}} d\varepsilon}{(2\pi)^{P/2}} \right| \\ & \leq \|\mathbf{A}^*\| \|\Pi\| L_1 \|\mathbf{S}\| \int_{\mathbb{R}^P} \|\varepsilon\|^2 \frac{e^{-\frac{\|\varepsilon\|^2}{2}} d\varepsilon}{(2\pi)^{P/2}} < \infty. \end{aligned} \quad (111)$$

Then, Equation (110) means that

$$\text{Tr}(\mathbf{S} \mathbf{A}^* \Pi \partial_{\mathbf{y}} \widehat{\mathbf{x}}(\mathbf{y}; \Lambda)) = \lim_{\nu \rightarrow 0} \mathbb{E}_\varepsilon \left\langle \mathbf{A}^* \Pi \frac{\widehat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \Lambda) - \widehat{\mathbf{x}}(\mathbf{y}; \Lambda)}{\nu}, \mathbf{S} \varepsilon \right\rangle. \quad (112)$$

Further, the majoration obtained in Equation (111) not depending on ζ (since L_1 does not depend on $\mathbf{y}$, as stated in Assumption 4), neither on ν , the limits on ν and the expected value with respect to Gaussian random noise ζ can be interchanged so that

$$\begin{aligned} \mathbb{E}_\zeta \text{Tr}(\mathbf{S} \mathbf{A}^* \Pi \partial_{\mathbf{y}} \widehat{\mathbf{x}}(\mathbf{y}; \Lambda)) &= \mathbb{E}_\zeta \lim_{\nu \rightarrow 0} \mathbb{E}_\varepsilon \left\langle \mathbf{A}^* \Pi \frac{\widehat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \Lambda) - \widehat{\mathbf{x}}(\mathbf{y}; \Lambda)}{\nu}, \mathbf{S} \varepsilon \right\rangle \\ &= \lim_{\nu \rightarrow 0} \mathbb{E}_{\zeta, \varepsilon} \left\langle \mathbf{A}^* \Pi \frac{\widehat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \Lambda) - \widehat{\mathbf{x}}(\mathbf{y}; \Lambda)}{\nu}, \mathbf{S} \varepsilon \right\rangle. \end{aligned} \quad (113)$$

giving the asymptotic unbiasedness of the Finite Difference Monte Carlo estimator of degrees of freedom and hence of the Finite Difference Monte Carlo SURE (17). $\square$

C Finite Difference Monte Carlo SUGAR

Proof. We remind that Finite Difference Monte Carlo SUGAR is composed of two terms, denoted $(\partial 1)$ and $(\partial 2)$ in the following:

$$\begin{aligned} \partial_{\Lambda} \hat{R}_{\nu, \epsilon}(\mathbf{y}; \Lambda | \mathcal{S}) &\triangleq \\ 2(\mathbf{A} \Phi \partial_{\Lambda} \hat{\mathbf{x}}(\mathbf{y}; \Lambda))^* \mathbf{A} (\Phi \hat{\mathbf{x}} - \mathbf{y}) &+ \frac{2}{\nu} \langle \mathbf{A}^* \mathbf{\Pi} (\partial_{\Lambda} \hat{\mathbf{x}}(\mathbf{y} + \nu \epsilon; \Lambda) - \partial_{\Lambda} \hat{\mathbf{x}}(\mathbf{y}; \Lambda)), \mathcal{S} \epsilon \rangle. \end{aligned} \quad \begin{aligned} (\partial 1) \qquad \qquad \qquad (\partial 2) \end{aligned} \quad (114)$$

Since the estimator $\hat{\mathbf{x}}(\mathbf{y}; \Lambda)$ is weakly differentiable with respect to Λ , so is the true risk $R[\hat{\mathbf{x}}](\Lambda)$. Thus, for any continuously differentiable test function $\varphi : \mathbb{R}^L \rightarrow \mathbb{R} \in C^1(\mathbb{V})$ with compact support denoted $\mathbb{V} \subset \mathbb{R}^L$, and any component $l \in \{1, \dots, L\}$ of the gradient of the risk $\partial_{\Lambda} R[\hat{\mathbf{x}}](\Lambda)$

$$\begin{aligned} \int_{\mathbb{R}^L} (\partial_{\Lambda} R[\hat{\mathbf{x}}](\Lambda))_l \varphi(\Lambda) d\Lambda &= \int_{\mathbb{V}} (\partial_{\Lambda} R[\hat{\mathbf{x}}](\Lambda))_l \varphi(\Lambda) d\Lambda \quad (115) \\ \text{(Weak differentiability)} &= - \int_{\mathbb{V}} R[\hat{\mathbf{x}}](\Lambda) (\partial_{\Lambda} \varphi(\Lambda))_l d\Lambda \\ \text{(Definition of the risk (10))} &= - \int_{\mathbb{V}} \mathbb{E}_{\zeta} \|\mathbf{\Pi} \hat{\mathbf{x}}(\mathbf{y}; \Lambda) - \mathbf{\Pi} \mathbf{x}\|_2^2 (\partial_{\Lambda} \varphi(\Lambda))_l d\Lambda \\ \text{(Theorem 1)} &= - \int_{\mathbb{V}} \mathbb{E}_{\zeta, \epsilon} \lim_{\nu \rightarrow 0} \hat{R}_{\nu, \epsilon}(\mathbf{y}; \Lambda | \mathcal{S}) (\partial_{\Lambda} \varphi(\Lambda))_l d\Lambda \\ \text{(Theorem 2)} &= - \int_{\mathbb{V}} \lim_{\nu \rightarrow 0} \mathbb{E}_{\zeta, \epsilon} \hat{R}_{\nu, \epsilon}(\mathbf{y}; \Lambda | \mathcal{S}) (\partial_{\Lambda} \varphi(\Lambda))_l d\Lambda \\ \text{(Dominated convergence)} &\stackrel{(\text{DC } 1)}{=} - \lim_{\nu \rightarrow 0} \int_{\mathbb{V}} \mathbb{E}_{\zeta, \epsilon} \hat{R}_{\nu, \epsilon}(\mathbf{y}; \Lambda | \mathcal{S}) (\partial_{\Lambda} \varphi(\Lambda))_l d\Lambda \\ \text{(Fubini)} &\stackrel{(\text{Fu } 1)}{=} - \lim_{\nu \rightarrow 0} \mathbb{E}_{\zeta, \epsilon} \int_{\mathbb{V}} \hat{R}_{\nu, \epsilon}(\mathbf{y}; \Lambda | \mathcal{S}) (\partial_{\Lambda} \varphi(\Lambda))_l d\Lambda \\ \text{(Proposition 1)} &= \lim_{\nu \rightarrow 0} \mathbb{E}_{\zeta, \epsilon} \int_{\mathbb{V}} \left(\partial_{\Lambda} \hat{R}_{\nu, \epsilon}(\mathbf{y}; \Lambda | \mathcal{S}) \right)_l \varphi(\Lambda) d\Lambda \\ &\stackrel{(\text{Fu } 2)}{=} \lim_{\nu \rightarrow 0} \int_{\mathbb{V}} \mathbb{E}_{\zeta, \epsilon} \left(\partial_{\Lambda} \hat{R}_{\nu, \epsilon}(\mathbf{y}; \Lambda | \mathcal{S}) \right)_l \varphi(\Lambda) d\Lambda \\ &\stackrel{(\text{DC } 2)}{=} \int_{\mathbb{V}} \lim_{\nu \rightarrow 0} \mathbb{E}_{\zeta, \epsilon} \left(\partial_{\Lambda} \hat{R}_{\nu, \epsilon}(\mathbf{y}; \Lambda | \mathcal{S}) \right)_l \varphi(\Lambda) d\Lambda. \end{aligned}$$

[\(DC 1\)](#) In order to apply dominated convergence theorem interchanging the limit on ν and the integral on $\mathbb{V}$, since φ is a test function with compact domain $\mathbb{V}$, we derive a bound of $\mathbb{E}_{\zeta, \epsilon} \hat{R}_{\nu, \epsilon}(\mathbf{y}; \Lambda | \mathcal{S})$ which is independent of both ν and Λ . Using the probability density functions we have

$$\mathbb{E}_{\zeta, \epsilon} \hat{R}_{\nu, \epsilon}(\mathbf{y}; \Lambda | \mathcal{S}) = \int_{\zeta} \int_{\epsilon} \hat{R}_{\nu, \epsilon}(\mathbf{y}; \Lambda | \mathcal{S}) \mathcal{G}_{\mathcal{S}}(\zeta) \mathcal{G}_I(\epsilon) d\zeta d\epsilon \quad (116)$$

where $\mathcal{G}_{\mathcal{S}}$ (resp. $\mathcal{G}_I$) denotes the Gaussian probability density function with covariance matrix $\mathcal{S}$ (resp. I)

$$\mathcal{G}_{\mathcal{S}}(\zeta) \triangleq \frac{\exp(-\|\zeta\|_{\mathcal{S}^{-1}}^2/2)}{\sqrt{(2\pi)^P |\det \mathcal{S}|}} \quad \left(\text{resp. } \mathcal{G}_I(\zeta) \triangleq \frac{\exp(-\|\zeta\|_2^2/2)}{\sqrt{(2\pi)^P}} \right). \quad (117)$$

We remind that $\widehat{R}_{\nu,\varepsilon}(\mathbf{y}; \mathbf{\Lambda}|\mathcal{S})$ is decomposed of three terms

$$\begin{aligned} \widehat{R}_{\nu,\varepsilon}(\mathbf{y}; \mathbf{\Lambda}|\mathcal{S}) &\triangleq \\ &\underbrace{\|\mathbf{A}(\Phi\widehat{\mathbf{x}} - \mathbf{y})\|_2^2}_{(1)} + \underbrace{\frac{1}{\nu} \langle \mathbf{A}^* \mathbf{\Pi} (\widehat{\mathbf{x}}(\mathbf{y} + \nu\varepsilon; \mathbf{\Lambda}|\mathcal{S}) - \widehat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}|\mathcal{S})), \mathcal{S}\varepsilon \rangle}_{(2)} - \underbrace{\text{Tr}(\mathbf{A}\mathcal{S}\mathbf{A}^*)}_{(3)}. \end{aligned} \quad (118)$$

which will be bounded separately.

(1) First, combining Assumptions (i) and (ii) of Assumption 4, we have

$$\|\widehat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}) - \widehat{\mathbf{x}}(\mathbf{0}_P; \mathbf{\Lambda})\| \leq L_1 \|\mathbf{y} - \mathbf{0}_P\| \implies \|\widehat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})\| \leq L_1 \|\mathbf{y}\|, \quad (119)$$

which can be used to bound first term (1) of (118) as

$$\begin{aligned} \|\mathbf{A}(\Phi\widehat{\mathbf{x}} - \mathbf{y})\| &\leq \|\mathbf{A}\| \|\Phi\widehat{\mathbf{x}} - \mathbf{y}\| \\ &\leq \|\mathbf{A}\| (\|\Phi\widehat{\mathbf{x}}\| + \|\mathbf{y}\|) \\ &\leq \|\mathbf{A}\| (\|\Phi\| \|\widehat{\mathbf{x}}\| + \|\mathbf{y}\|) \\ &\leq \|\mathbf{A}\| (\|\Phi\| L_1 \|\mathbf{y}\| + \|\mathbf{y}\|) \\ &\leq \|\mathbf{A}\| (\|\Phi\| L_1 + 1) \|\mathbf{y}\|. \end{aligned} \quad (120)$$

Since by definition $\zeta = \mathbf{y} - \Phi\mathbf{x}$, $\mathbf{y} \mapsto \|\mathbf{y}\|$ is integrable against the Gaussian density $\mathcal{G}_{\mathcal{S}}(\zeta)$, and the above domination being independent of ν , it enable us to apply dominated convergence.

(2) Making use of the domination of Equation (109) we have

$$\begin{aligned} &\left| \left\langle \mathbf{A}^* \mathbf{\Pi} \frac{(\widehat{\mathbf{x}}(\mathbf{y} + \nu\varepsilon; \mathbf{\Lambda}) - \widehat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda}))}{\nu}, \mathcal{S}\varepsilon \right\rangle \right| \\ &\leq \|\mathbf{A}^*\| \|\mathbf{\Pi}\| L_1 \|\varepsilon\| \|\mathcal{S}\| \|\varepsilon\|, \end{aligned} \quad (121)$$

and $\|\varepsilon\|^2$ being integrable against $\mathcal{G}_I(\varepsilon)$, dominated convergence applied.

(3) The third term being constant, the domination is obvious.

Putting altogether the majoration of (1), (2) and (3)

$$\begin{aligned} \left| \mathbb{E}_{\zeta,\varepsilon} \widehat{R}_{\nu,\varepsilon}(\mathbf{y}; \mathbf{\Lambda}|\mathcal{S}) \right| &\leq \int_{\zeta} \int_{\varepsilon} \|\mathbf{A}\| (\|\Phi\| L_1 + 1) \|\mathbf{y}\| \mathcal{G}_{\mathcal{S}}(\zeta) \mathcal{G}_I(\varepsilon) d\zeta d\varepsilon \\ &\quad + \int_{\zeta} \int_{\varepsilon} \|\mathbf{A}^*\| \|\mathbf{\Pi}\| L_1 \|\mathcal{S}\| \|\varepsilon\|^2 \mathcal{G}_{\mathcal{S}}(\zeta) \mathcal{G}_I(\varepsilon) d\zeta d\varepsilon \\ &\quad + \int_{\zeta} \int_{\varepsilon} \text{Tr}(\mathbf{A}\mathcal{S}\mathbf{A}^*) \mathcal{G}_{\mathcal{S}}(\zeta) \mathcal{G}_I(\varepsilon) d\zeta d\varepsilon \leq \infty, \end{aligned} \quad (122)$$

the majoration being independent of ν and $\mathbf{\Lambda}$ dominated convergence applies.

(Fu 1) The above domination of Equation (122) being independent of $\mathbf{\Lambda}$, then Fubini's theorem applies.

(Fu 2) The first term of (114), denoted as $(\partial 1)$ can be easily dominated by an integrable function. Indeed, Assumption 5 implies that $\partial_{\mathbf{\Lambda}} \widehat{\mathbf{x}}(\mathbf{y}; \mathbf{\Lambda})$ is uniformly

bounded by L_2 , independently of $\mathbf{y}$. Then it follows from Cauchy-Schwarz inequality and the domination of Equation (120)

$$\begin{aligned} 2 \|(\mathbf{A} \Phi \partial_{\Lambda} \hat{\mathbf{x}}(\mathbf{y}; \Lambda))^* \mathbf{A} (\Phi \hat{\mathbf{x}} - \mathbf{y})\| &\leq 2 \|\mathbf{A} \Phi \partial_{\Lambda} \hat{\mathbf{x}}(\mathbf{y}; \Lambda)\| \|\mathbf{A} (\Phi \hat{\mathbf{x}} - \mathbf{y})\|, \\ &\leq 2 \|\mathbf{A}\| \|\Phi\| L_2 \|\mathbf{A}\| (\Phi L_1 + 1) \|\mathbf{y}\|. \end{aligned} \quad (123)$$

Hence, since $\|\mathbf{y}\|$ is integrable against $\mathcal{G}_{\mathcal{S}}(\mathbf{y} - \Phi \mathbf{x}) \mathcal{G}_I(\varepsilon)$ and the domination being independent of ν , both Fubini and dominated convergence theorems apply.

The second term of (114), denoted as $(\partial 2)$, corresponding to the derivative of the estimation of degrees of freedom, can be rewritten as

$$\frac{2}{\nu} \langle \mathbf{A}^* \Pi (\partial_{\Lambda} \hat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \Lambda) - \partial_{\Lambda} \hat{\mathbf{x}}(\mathbf{y}; \Lambda)), \mathcal{S} \varepsilon \rangle \triangleq \frac{2}{\nu} \langle u(\zeta + \nu \varepsilon; \Lambda) - u(\zeta; \Lambda), \varepsilon \rangle$$

where we set $u(\mathbf{z}; \Lambda) \triangleq \mathcal{S} \mathbf{A}^* \Pi \partial_{\Lambda} \hat{\mathbf{x}}(\Phi \mathbf{x} + \mathbf{z}; \Lambda)$. Since $\partial_{\Lambda} \hat{\mathbf{x}}(\mathbf{y}; \Lambda)$ is uniformly bounded by L_2 , independently of $\mathbf{y}$, and all the linear operators are assumed to be bounded, then $\Lambda \mapsto u(\mathbf{z}; \Lambda)$ is bounded by some $L_u > 0$, independently of $\mathbf{z}$. Then

$$\begin{aligned} &\mathbb{E}_{\zeta, \varepsilon} \left[\frac{2}{\nu} \langle \mathbf{A}^* \Pi (\partial_{\Lambda} \hat{\mathbf{x}}(\mathbf{y} + \nu \varepsilon; \Lambda) - \partial_{\Lambda} \hat{\mathbf{x}}(\mathbf{y}; \Lambda)), \mathcal{S} \varepsilon \rangle \right] \\ &= \int_{\zeta} \int_{\varepsilon} \frac{2}{\nu} \langle u(\zeta + \nu \varepsilon; \Lambda) - u(\zeta; \Lambda), \varepsilon \rangle \mathcal{G}_{\mathcal{S}}(\zeta) \mathcal{G}_I(\varepsilon) d\zeta d\varepsilon \\ &= \int_{\zeta} \int_{\varepsilon} \frac{2}{\nu} \langle u(\zeta + \nu \varepsilon; \Lambda), \varepsilon \rangle \mathcal{G}_{\mathcal{S}}(\zeta) \mathcal{G}_I(\varepsilon) d\zeta d\varepsilon - \int_{\zeta} \int_{\varepsilon} \frac{2}{\nu} \langle u(\zeta; \Lambda), \varepsilon \rangle \mathcal{G}_{\mathcal{S}}(\zeta) \mathcal{G}_I(\varepsilon) d\zeta d\varepsilon \\ &= \int_{\zeta} \int_{\varepsilon} \frac{2}{\nu} \langle u(\zeta; \Lambda), \varepsilon \rangle (\mathcal{G}_{\mathcal{S}}(\zeta - \nu \varepsilon) - \mathcal{G}_{\mathcal{S}}(\zeta)) \mathcal{G}_I(\varepsilon) d\zeta d\varepsilon. \end{aligned} \quad (124)$$

Further, the following majoration holds

$$\left\| \frac{2}{\nu} \langle u(\zeta; \Lambda), \varepsilon \rangle (\mathcal{G}_{\mathcal{S}}(\zeta - \nu \varepsilon) - \mathcal{G}_{\mathcal{S}}(\zeta)) \mathcal{G}_I(\varepsilon) \right\| \leq \frac{2}{\nu} L_u \|\varepsilon\| |\mathcal{G}_{\mathcal{S}}(\zeta - \nu \varepsilon) - \mathcal{G}_{\mathcal{S}}(\zeta)| \mathcal{G}_I(\varepsilon). \quad (125)$$

Up to a unitary variable change (see Appendix A, with $\vartheta = \mathbf{V}^{-1} \varepsilon$), we can assume that the covariance matrix is diagonal, with diagonal terms $(s_i^2)_{i=1}^P$ so that

$$\mathcal{G}_{\mathcal{S}}(\zeta) = \prod_{i=1}^P \frac{\exp(-|\zeta_i|^2 / 2s_i^2)}{\sqrt{2\pi s_i^2}} \quad (126)$$

and we define the one-dimensional Gaussian densities as

$$g_{s_i^2}(\zeta_i) \triangleq \frac{\exp(-|\zeta_i|^2 / 2s_i^2)}{\sqrt{2\pi s_i^2}}. \quad (127)$$

From Taylor inequality,

$$|g_{s_i^2}(\zeta_i - \nu \varepsilon_i) - g_{s_i^2}(\zeta_i)| \leq \int_{(0, \nu \varepsilon_i)} |g'_{s_i^2}(\zeta_i - \tau)| d\tau, \quad (128)$$

where $(0, \nu\varepsilon_i)$ denotes the ordered interval, taking into account that ε_i might be negative that is

$$(0, \nu\varepsilon_i) = \begin{cases} [0, \nu\varepsilon_i] & \text{if } \varepsilon_i \geq 0 \\ [\nu\varepsilon_i, 0] & \text{else} \end{cases} \quad (129)$$

then

$$\begin{aligned} \int_{\zeta_i} |g_{s_i^2}(\zeta_i - \nu\varepsilon_i) - g_{s_i^2}(\zeta_i)| d\zeta_i &\leq \int_{\zeta_i} \int_{(0, \nu\varepsilon_i)} |g'_{s_i^2}(\zeta_i - \tau)| d\tau d\zeta_i \\ &\leq \int_{(0, \nu\varepsilon_i)} \left(\int_{\zeta_i} |g'_{s_i^2}(\zeta_i - \tau)| d\zeta_i \right) d\tau \\ &= \left(\int_{\mathbb{R}} |g'_{s_i^2}(t)| dt \right) \nu|\varepsilon_i| < +\infty \end{aligned}$$

since the derivative of the Gaussian density is integrable over $\mathbb{R}$.

Going back to the integrals over variables $\zeta, \varepsilon \in \mathbb{R}^P$ of Equation (124)

$$\begin{aligned} &\left\| \int_{\zeta} \int_{\varepsilon} \frac{2}{\nu} \langle u(\zeta; \Lambda), \varepsilon \rangle (\mathcal{G}_{\mathcal{S}}(\zeta - \nu\varepsilon) - \mathcal{G}_{\mathcal{S}}(\zeta)) \mathcal{G}_{\mathbf{I}}(\varepsilon) d\zeta d\varepsilon \right\| \quad (130) \\ &\leq \int_{\zeta} \int_{\varepsilon} \frac{2}{\nu} L_u \|\varepsilon\| |\mathcal{G}_{\mathcal{S}}(\zeta - \nu\varepsilon) - \mathcal{G}_{\mathcal{S}}(\zeta)| \mathcal{G}_{\mathbf{I}}(\varepsilon) d\zeta d\varepsilon \\ &= \int_{\varepsilon} \frac{2}{\nu} L_u \|\varepsilon\| \prod_{i=1}^P \left(\int_{\zeta_i} |g_{s_i^2}(\zeta_i - \nu\varepsilon_i) - g_{s_i^2}(\zeta_i)| d\zeta_i \right) \mathcal{G}_{\mathbf{I}}(\varepsilon) d\varepsilon \\ &\leq \int_{\varepsilon} \frac{2}{\nu} L_u \|\varepsilon\| \prod_{i=1}^P \left(\left(\int_{\mathbb{R}} |g'_{s_i^2}(t)| dt \right) \nu|\varepsilon_i| \right) \mathcal{G}_{\mathbf{I}}(\varepsilon) d\varepsilon \\ |\varepsilon_i| \leq \|\varepsilon\| &\leq \int_{\varepsilon} 2\nu^{P-1} \|L_u \varepsilon\| \|\varepsilon\|^P \prod_{i=1}^P \left(\int_{\mathbb{R}} |g'_{s_i^2}(t)| dt \right) \mathcal{G}_{\mathbf{I}}(\varepsilon) d\varepsilon \\ (0 < \nu \leq 1) &\leq \int_{\varepsilon} 2L_u \|\varepsilon\|^{P+1} \prod_{i=1}^P \left(\int_{\mathbb{R}} |g'_{s_i^2}(t)| dt \right) \mathcal{G}_{\mathbf{I}}(\varepsilon) d\varepsilon < +\infty \end{aligned}$$

Indeed, $\|\cdot\|$ being the euclidean norm ($\forall i$) $|\varepsilon_i| \leq \|\varepsilon\|$. Moreover, since we are interested in the limit $\nu \rightarrow 0$, we can assume without loss of generality that $0 < \nu \leq 1$ and thus $\nu^{P-1} \leq 1$. We conclude using the fact that any power of $\|\varepsilon\|$ is integrable against $\mathcal{G}_{\mathbf{I}}(\varepsilon)$, combined to the fact that the support $\mathbb{V}$ of φ is compact, which enable to apply Fubini's theorem to exchange $\int_{\mathbb{V}}$ and $\mathbb{E}_{\zeta, \varepsilon}$.

(DC 2) The above majoration does not depends on ν . Further, the Lipschitzianity assumptions provides the existence P-*a.s.* of

$$\lim_{\nu \rightarrow 0} \mathbb{E}_{\zeta, \varepsilon} \left(\partial_{\Lambda} \widehat{R}_{\nu, \varepsilon}(\mathbf{y}; \Lambda | \mathcal{S}) \right)_l$$

then dominated convergence theorem applies to invert $\lim_{\nu \rightarrow 0}$ and $\int_{\mathbb{V}}$ which completes the proof. $\square$

D Constant term of Stein Unbiased Risk Estimate

Proof. Because of the cyclicity of the trace, one has $\text{Tr}(\mathbf{A}\mathbf{S}\mathbf{A}^*) = \text{Tr}(\mathbf{A}^*\mathbf{A}\mathbf{S})$, then using the definition of $\mathbf{A} \triangleq \mathbf{\Pi}(\mathbf{\Phi}^*\mathbf{\Phi})^{-1}\mathbf{\Phi}^*$,

$$\mathbf{A}^*\mathbf{A} = \mathbf{\Phi}(\mathbf{\Phi}^*\mathbf{\Phi})^{-1}\mathbf{\Pi}^*\mathbf{\Pi}(\mathbf{\Phi}^*\mathbf{\Phi})^{-1}\mathbf{\Phi}^*$$

then turning to the block-matrix form to perform the products

$$\begin{aligned} & (\mathbf{\Phi}^*\mathbf{\Phi})^{-1}\mathbf{\Pi}^*\mathbf{\Pi}(\mathbf{\Phi}^*\mathbf{\Phi})^{-1} \\ &= \frac{1}{(F_0F_2 - F_1^2)^2} \begin{pmatrix} F_0\mathbf{I}_{N/2} & -F_1\mathbf{I}_{N/2} \\ -F_1\mathbf{I}_{N/2} & F_2\mathbf{I}_{N/2} \end{pmatrix} \begin{pmatrix} \mathbf{I}_{N/2} & \mathbf{Z}_{N/2} \\ \mathbf{Z}_{N/2} & \mathbf{I}_{N/2} \end{pmatrix} \begin{pmatrix} F_0\mathbf{I}_{N/2} & -F_1\mathbf{I}_{N/2} \\ -F_1\mathbf{I}_{N/2} & F_2\mathbf{I}_{N/2} \end{pmatrix} \\ &= \frac{1}{(F_0F_2 - F_1^2)^2} \begin{pmatrix} F_0\mathbf{I}_{N/2} & -F_1\mathbf{I}_{N/2} \\ -F_1\mathbf{I}_{N/2} & F_2\mathbf{I}_{N/2} \end{pmatrix} \begin{pmatrix} F_0\mathbf{Z}_{N/2} & -F_1\mathbf{I}_{N/2} \\ \mathbf{Z}_{N/2} & \mathbf{I}_{N/2} \end{pmatrix} \\ &= \frac{1}{(F_0F_2 - F_1^2)^2} \begin{pmatrix} F_0^2\mathbf{I}_{N/2} & -F_0F_1\mathbf{I}_{N/2} \\ -F_0F_1\mathbf{I}_{N/2} & F_1^2\mathbf{I}_{N/2} \end{pmatrix}. \end{aligned}$$

Using again cyclicity of the trace

$$\begin{aligned} \text{Tr}(\mathbf{A}\mathbf{S}\mathbf{A}^*) &= \text{Tr} \left(\mathbf{\Phi}(\mathbf{\Phi}^*\mathbf{\Phi})^{-1}\mathbf{\Pi}^*\mathbf{\Pi}(\mathbf{\Phi}^*\mathbf{\Phi})^{-1}\mathbf{\Phi}^*\mathbf{S} \right) \\ &= \text{Tr} \left(\mathbf{\Phi} \frac{1}{(F_0F_2 - F_1^2)^2} \begin{pmatrix} F_0^2\mathbf{I}_{N/2} & -F_0F_1\mathbf{I}_{N/2} \\ -F_0F_1\mathbf{I}_{N/2} & F_1^2\mathbf{I}_{N/2} \end{pmatrix} \mathbf{\Phi}^*\mathbf{S} \right) \\ &= \text{Tr} \left(\frac{1}{(F_0F_2 - F_1^2)^2} \begin{pmatrix} F_0^2\mathbf{I}_{N/2} & -F_0F_1\mathbf{I}_{N/2} \\ -F_0F_1\mathbf{I}_{N/2} & F_1^2\mathbf{I}_{N/2} \end{pmatrix} \mathbf{\Phi}^*\mathbf{S}\mathbf{\Phi} \right) \end{aligned}$$

Then using the action of $\mathbf{\Phi}$ and $\mathbf{\Phi}^*$, explicited in Formula (66) we have the matrix representation

$$\mathbf{\Phi} = \begin{pmatrix} \mathbf{I}_{N/2} & \mathbf{I}_{N/2} \\ 2\mathbf{I}_{N/2} & \mathbf{I}_{N/2} \\ \vdots & \vdots \\ J\mathbf{I}_{N/2} & \mathbf{I}_{N/2} \end{pmatrix} \in \mathbb{R}^{JN_1N_2 \times 2N_1N_2}. \quad (131)$$

Using of the decomposition of $\mathbf{S}$ into J^2 blocks $\mathbf{S}_j^{j'} = \mathbf{C}_j^{j'} \mathbf{\Xi}_j^{j'} \in \mathbb{R}^{N/2 \times N/2}$, we obtain

$$\mathbf{\Phi}^*\mathbf{S}\mathbf{\Phi} = \begin{pmatrix} \sum_{j,j'} jj' \mathbf{S}_j^{j'} & \sum_{j,j'} j' \mathbf{S}_j^{j'} \\ \sum_{j,j'} j \mathbf{S}_j^{j'} & \sum_{j,j'} \mathbf{S}_j^{j'} \end{pmatrix}, \quad 1 \leq j, j' \leq J. \quad (132)$$

It follows

$$\begin{aligned}
& \text{Tr}(\mathbf{A}\mathbf{S}\mathbf{A}^*) \\
&= \text{Tr} \left(\frac{1}{(F_0F_2 - F_1^2)^2} \begin{pmatrix} F_0^2 \mathbf{I}_{N/2} & -F_0F_1 \mathbf{I}_{N/2} \\ -F_0F_1 \mathbf{I}_{N/2} & F_1^2 \mathbf{I}_{N/2} \end{pmatrix} \begin{pmatrix} \sum_{j,j'} jj' \mathbf{S}_j^{j'} & \sum_{j,j'} j' \mathbf{S}_j^{j'} \\ \sum_{j,j'} j \mathbf{S}_j^{j'} & \sum_{j,j'} \mathbf{S}_j^{j'} \end{pmatrix} \right) \\
&= \frac{1}{(F_0F_2 - F_1^2)^2} \text{Tr} \begin{pmatrix} \sum_{j,j'} F_0^2 jj' \mathbf{S}_j^{j'} - F_0F_1 j' \mathbf{S}_j^{j'} & \sum_{j,j'} F_0^2 j' \mathbf{S}_j^{j'} - F_0F_1 \mathbf{S}_j^{j'} \\ \sum_{j,j'} F_1^2 j \mathbf{S}_j^{j'} - F_0F_1 jj' \mathbf{S}_j^{j'} & \sum_{j,j'} F_1^2 \mathbf{S}_j^{j'} - F_0F_1 j' \mathbf{S}_j^{j'} \end{pmatrix} \\
&= \frac{1}{(F_0F_2 - F_1^2)^2} \text{Tr} \left(\sum_{j,j'} \left(F_0^2 jj' \mathbf{S}_j^{j'} - F_0F_1 j' \mathbf{S}_j^{j'} + F_1^2 \mathbf{S}_j^{j'} - F_0F_1 j' \mathbf{S}_j^{j'} \right) \right). \tag{133}
\end{aligned}$$

Then, one can remark that

$$\text{Tr}(\mathbf{S}_j^{j'}) \triangleq \sum_{\underline{n} \in \Omega} \mathbf{S}_{j,\underline{n}}^{j'} = \sum_{\underline{n} \in \Omega} \mathcal{C}_j^{j'} = \frac{N}{2} \mathcal{C}_j^{j'}. \tag{134}$$

since the filter $\Xi_j^{j'}$ is supposed to be normalized, in the sense that its maximum value equals 1. Finally

$$\text{Tr}(\mathbf{A}\mathbf{S}\mathbf{A}^*) = \frac{N/2}{(F_0F_2 - F_1^2)^2} \sum_{j,j'} \left(F_0^2 jj' \mathcal{C}_j^{j'} - F_0F_1 j' \mathcal{C}_j^{j'} + F_1^2 \mathcal{C}_j^{j'} - F_0F_1 j' \mathcal{C}_j^{j'} \right). \tag{135}$$

□

References

- [1] A. C. Aitkin. On least squares and linear combination of observations. *Proceedings of the Royal Society of Edinburgh*, 55:42–48, 1935.
- [2] R. Ammanouil, A. Ferrari, D. Mary, C. Ferrari, and F. Loi. A parallel and automatically tuned algorithm for multispectral image deconvolution. *Monthly Notices of the Royal Astronomical Society*, 490(1):37–49, 2019.
- [3] J.-F. Aujol, G. Gilboa, T.F. Chan, and S. Osher. Structure-texture image decomposition—modeling, algorithms, and parameter selection. *Int. J. Comp. Vis.*, 67(1):111–136, 2006.
- [4] H. H. Bauschke and P. L. Combettes. *Convex analysis and monotone operator theory in Hilbert spaces*, volume 408. Springer, 2011.
- [5] A. Beck and M. Teboulle. A fast iterative shrinkage-thresholding algorithm for linear inverse problems. *SIAM J. Imaging Sci.*, 2(1):183–202, 2009.
- [6] J. Berger. Minimax estimation of a multivariate normal mean under arbitrary quadratic loss. *J. Multivariate Anal.*, 6(2):256–264, 1976.
- [7] J. Bergstra and Y. Bengio. Random search for hyper-parameter optimization. *Journal of Machine Learning Research*, 13(Feb):281–305, 2012.

- [8] J. Bergstra, D. Yamins, and D. D. Cox. Making a science of model search: Hyperparameter optimization in hundreds of dimensions for vision architectures. Atlanta, USA, June 16–21 2013. Jmlr.
- [9] Q. Bertrand, Q. Klopfenstein, M. Blondel, S. Vaiter, A. Gramfort, and J. Salmon. Implicit differentiation of Lasso-type models for hyperparameter optimization, 2020.
- [10] Å. Björck. *Numerical methods for least squares problems*. SIAM, 1996.
- [11] T. Blu and F. Luisier. The SURE-LET approach to image denoising. *IEEE Trans. Image Process.*, 16(11):2778–2786, 2007.
- [12] R. H. Byrd, P. Lu, J. Nocedal, and C. Zhu. A limited memory algorithm for bound constrained optimization. *J. Sci. Comput.*, 16(5):1190–1208, 1995.
- [13] J.-F. Cai, B. Dong, S. Osher, and Z. Shen. Image restoration: Total variation, wavelet frames, and beyond. *J. Amer. Math. Soc.*, 25:1033–1089, May 2012.
- [14] X. Cai, R. Chan, C.-B. Schonlieb, G. Steidl, and T. Zeng. Linkage between piecewise constant mumford-shah model and rof model and its virtue in image segmentation. *arXiv preprint arXiv:1807.10194*, 2018.
- [15] X. Cai and G. Steidl. Multiclass segmentation by iterated rof thresholding. In *Int. Workshop on Energy Minimization Methods in Comp. Vis. and Pat. Rec.*, pages 237–250. Springer, 2013.
- [16] A. Chambolle and T. Pock. A first-order primal-dual algorithm for convex problems with applications to imaging. *J. Math. Imag. Vis.*, 40(1):120–145, 2011.
- [17] T.F. Chan and J.J. Shen. Variational image inpainting. *Comm. Pure Applied Math.*, 58(5), May 2005.
- [18] C. Chaux and L. Blanc-Féraud. Estimation d’hyperparamètres pour la résolution de problèmes inverses à l’aide d’ondelettes. In *XXIIe colloque GRETSI (Traitement du Signal et des Images)*, Dijon, France, Sept. 8–11 2009. GRETSI, Groupe d’Etudes du Traitement du Signal et des Images.
- [19] C. Chaux, L. Duval, A. Benazza-Benyahia, and J.-C. Pesquet. A nonlinear stein-based estimator for multichannel image denoising. *IEEE Trans. Signal Process.*, 56(8):3855–3870, 2008.
- [20] P. L. Combettes and J.-C. Pesquet. Proximal splitting methods in signal processing. In H. H. Bauschke, R. S. Burachik, P. L. Combettes, V. Elser, D. R. Luke, and H. Wolkowicz, editors, *Fixed-Point Algorithms for Inverse Problems in Science and Engineering*, pages 185–212. Springer-Verlag, New York, 2011.
- [21] L. Condat, D. Kitahara, A. Contreras, and A. Hirabayashi. Proximal splitting algorithms: Relax them all!, 2019.

- [22] F. E. Curtis, T. Mitchell, and M. L. Overton. A BFGS-SQP method for nonsmooth, nonconvex, constrained optimization and its evaluation using relative minimization profiles. *Optim. Methods Softw.*, 32(1):148–181, 2017.
- [23] C.-A. Deledalle, F. Tupin, and L. Denis. Poisson NL means: Unsupervised non local means for Poisson noise. In *Proc. Int. Conf. Image Process.*, pages 801–804. IEEE, 2010.
- [24] C.-A. Deledalle, S. Vaiter, J. Fadili, and G. Peyré. Stein unbiased gradient estimator of the risk (SUGAR) for multiple parameter selection. *SIAM J. Imaging Sci.*, 7(4):2448–2487, 2014.
- [25] L. Desbat and D. Girard. The “minimum reconstruction error” choice of regularization parameters: Some more efficient methods and their application to deconvolution problems. *J. Sci. Comput.*, 16(6):1387–1403, 1995.
- [26] D. L. Donoho and I. M. Johnstone. Adapting to unknown smoothness via wavelet shrinkage. *J. American Statist. Ass.*, 90(432):1200–1224, 1995.
- [27] D. L. Donoho and J. M. Johnstone. Ideal spatial adaptation by wavelet shrinkage. *biometrika*, 81(3):425–455, 1994.
- [28] C. Dossal, M. Kachour, J. Fadili, G. Peyré, and C. Chesneau. The degrees of freedom of the Lasso for general design matrix. *Statistica Sinica*, pages 809–828, 2013.
- [29] B. Efron. How biased is the apparent error rate of a prediction rule? *J. Am. Stat. Assoc.*, 81(394):461–470, 1986.
- [30] Yonina C Eldar. Generalized SURE for exponential families: Applications to regularization. *IEEE Trans. Signal Process.*, 57(2):471–481, 2008.
- [31] L. Elden. Algorithms for the regularization of ill-conditioned least squares problems. *BIT Numer. Math.*, 17(2):134–145, 1977.
- [32] L. C. Evans and R. F. Gariepy. *Measure theory and fine properties of functions*. Chapman and Hall/CRC, 2015.
- [33] N. P. Galatsanos and A. K. Katsaggelos. Methods for choosing the regularization parameter and estimating the noise variance in image restoration and their relation. *IEEE Trans. Image Process.*, 1(3):322–336, 1992.
- [34] A. Girard. A fast Monte-Carlo cross-validation procedure for large least squares problems with noisy data. *Numerische Mathematik*, 56(1):1–23, 1989.
- [35] G. H. Golub, P. C. Hansen, and D. P. O’Leary. Tikhonov regularization and total least squares. *SIAM J. Matrix Anal. Appl.*, 21(1):185–194, 1999.
- [36] G. H. Golub, M. Heath, and G. Wahba. Generalized cross-validation as a method for choosing a good ridge parameter. *Technometrics*, 21(2):215–223, 1979.
- [37] P. C. Hansen and D. P. O’Leary. The use of the L-curve in the regularization of discrete ill-posed problems. *J. Sci. Comput.*, 14(6):1487–1503, 1993.

- [38] H. M. Hudson. A natural identity for exponential families with applications in multiparameter estimation. *Ann. Stat.*, 6(3):473–484, 1978.
- [39] H. M. Hudson and T. C. M. Lee. Maximum likelihood restoration and choice of smoothing parameter in deconvolution of image data subject to Poisson noise. *Computational statistics & data analysis*, 26(4):393–410, 1998.
- [40] K. Kato. On the degrees of freedom in shrinkage estimation. *J. Multivariate Anal.*, 100(7):1338–1352, 2009.
- [41] C. L. Lawson and R. J. Hanson. *Solving least squares problems*, volume 15. SIAM, 1995.
- [42] Y. Le Montagner, E. D. Angelini, and J.-C. Olivo-Marin. An unbiased risk estimator for image denoising in the presence of mixed poisson–gaussian noise. *IEEE Trans. Image Process.*, 23(3):1255–1268, 2014.
- [43] K. C. Li. From stein’s unbiased risk estimates to the method of generalized cross validation. *Ann. Stat.*, pages 1352–1377, 1985.
- [44] F. Luisier, T. Blu, and M. Unser. A new SURE approach to image denoising: Interscale orthonormal wavelet thresholding. *IEEE Trans. Image Process.*, 16(3):593–606, 2007.
- [45] F. Luisier, T. Blu, and M. Unser. Image denoising in mixed poisson–gaussian noise. *IEEE Trans. Image Process.*, 20(3):696–708, 2010.
- [46] S. Mallat. *A Wavelet Tour of Signal Processing, Third Edition: The Sparse Way*. Academic Press, Inc., Orlando, FL, USA, 3rd edition, 2008.
- [47] M. Meyer and M. Woodroffe. On the degrees of freedom in shape-restricted regression. *Ann. Stat.*, pages 1083–1104, 2000.
- [48] J. D. B. Nelson, C. Naornita, and A. Isar. Semi-local scaling exponent estimation with box-penalty constraints and total-variation regularization. *IEEE Trans. Image Process.*, 25(7):3167–3181, 2016.
- [49] J. Nocedal and S. Wright. *Numerical optimization*. Springer Science & Business Media, 2006.
- [50] N. Parikh and S. Boyd. Proximal algorithms. *Foundations and Trends® in Optimization*, 1(3):127–239, 2014.
- [51] B. Pascal, N. Pustelnik, and P. Abry. Nonsmooth Convex Joint Estimation of Local Regularity and Local Variance for Fractal Texture Segmentation. *arXiv e-prints*, page arXiv:1910.05246, Oct 2019.
- [52] Barbara Pascal, Nelly Pustelnik, Patrice Abry, Marion Serres, and Valérie Vidal. Joint estimation of local variance and local regularity for texture segmentation. application to multiphase flow characterization. In *Proc. Int. Conf. Image Process.*, pages 2092–2096. IEEE, 2018.
- [53] J.-C. Pesquet, A. Benazza-Benyahia, and C. Chaux. A SURE approach for digital signal/image deconvolution problems. *IEEE Trans. Signal Process.*, 57(12):4616–4632, 2009.

- [54] O. Pont, A. Turiel, and H. Yahia. An optimized algorithm for the evaluation of local singularity exponents in digital signals. In *International Workshop on Combinatorial Image Analysis*, pages 346–357. Springer, 2011.
- [55] N. Pustelnik, A. Benazza-Benhayia, Y. Zheng, and J.-C. Pesquet. Wavelet-based image deconvolution and reconstruction. *Wiley Encyclopedia of Electrical and Electronics Engineering*, Feb. 2016.
- [56] N. Pustelnik, H. Wendt, P. Abry, and N. Dobigeon. Combining local regularity estimation and total variation optimization for scale-free texture segmentation. *IEEE Trans. Computational Imaging*, 2(4):468–479, 2016.
- [57] S. Ramani, T. Blu, and M. Unser. Monte-carlo SURE: A black-box optimization of regularization parameters for general denoising algorithms. *IEEE Trans. Image Process.*, 17(9):1540–1554, 2008.
- [58] R. T. Rockafellar. *Convex analysis*. Number 28. Princeton university press, 1970.
- [59] L. I. Rudin, S. Osher, and E. Fatemi. Nonlinear total variation based noise removal algorithms. *Physica D: nonlinear phenomena*, 60(1-4):259–268, 1992.
- [60] C. M. Stein. Estimation of the mean of a multivariate normal distribution. *Ann. Stat.*, pages 1135–1151, 1981.
- [61] T. Strutz. *Data fitting and uncertainty: A practical introduction to weighted least squares and beyond*. Vieweg and Teubner, 2010.
- [62] A. M. Thompson, J. C. Brown, J. W. Kay, and D. M. Titterton. A study of methods of choosing the smoothing parameter in image restoration by regularization. *IEEE Trans. Pattern Anal. Match. Int.*, (4):326–339, 1991.
- [63] R. Tibshirani. Regression shrinkage and selection via the lasso. *Journal of the Royal Statistical Society: Series B (Methodological)*, 58(1):267–288, 1996.
- [64] R. Tibshirani and L. Wasserman. Steins unbiased risk estimate. *Course notes from “Statistical Machine Learning”*, pages 1–12, 2015.
- [65] R. J. Tibshirani and J. Taylor. Degrees of freedom in lasso problems. *Ann. Stat.*, 40(2):1198–1232, 2012.
- [66] A. Tikhonov. Tikhonov regularization of incorrectly posed problems. *Soviet Mathematics Doklady*, 4:1624–1627, 1963.
- [67] S. Vaiter, C. Deledalle, J. Fadili, G. Peyré, and C. Dossal. The degrees of freedom of partly smooth regularizers. *Ann. Inst. Stat. Math*, 69(4):791–832, 2017.
- [68] Darryl Veitch and Patrice Abry. A wavelet-based joint estimator of the parameters of long-range dependence. *IEEE Trans. Inform. Theory*, 45(3):878–897, 1999.

- [69] C. Vonesch, S. Ramani, and M. Unser. Recursive risk estimation for non-linear image deconvolution with a wavelet-domain sparsity constraint. In *Proc. Int. Conf. Image Process.*, pages 665–668. IEEE, 2008.
- [70] H. Wendt. *Contributions of Wavelet Leaders and Bootstrap to Multifractal Analysis: Images, Estimation Performance, Dependence Structure and Vanishing Moments. Confidence Intervals and Hypothesis Tests*. PhD thesis, Ecole Normale Supérieure de Lyon, 2008.
- [71] H. Wendt, P. Abry, S. Jaffard, H. Ji, and Z. Shen. Wavelet leader multifractal analysis for texture classification. In *Proc. Int. Conf. Image Process.*, pages 3829–3832. IEEE, 2009.
- [72] H. Wendt, S. G. Roux, P. Abry, and S. Jaffard. Wavelet leaders and bootstrap for multifractal analysis of images. *Signal Process.*, 89(6):1100–1114, 2009.
- [73] X. Xie, S.C. Kou, and L. D. Brown. Sure estimates for a heteroscedastic hierarchical model. *J. American Statist. Ass.*, 107(500):1465–1479, 2012.