

SMOS IMAGES RESTORATION FROM L1A DATA: A SPARSITY-BASED VARIATIONAL APPROACH

J. Preciozzi, P. Musé

Univ. de la República, Uruguay

A. Almansa[†], S. Durand[‡]

[†] LTCI - Telecom ParisTech
[‡] MAP5 - Univ. René Descartes

A. Khazaal, B. Rougé

CESBIO
31400 Toulouse, France

ABSTRACT

Data degradation by radio frequency interferences (RFI) is one of the major challenges that SMOS and other interferometers radiometers missions have to face. Although a great number of the illegal emitters were turned off since the mission was launched, not all of the sources were completely removed. Moreover, the data obtained previously is already corrupted by these RFI. Thus, the recovery of brightness temperature from corrupted data by image restoration techniques is of major interest. In this work we propose a variational approach to recover a super-resolved, denoised brightness temperature map based on two spatial components: an image u that models the brightness temperature and an image o modeling the RFI. The approach is totally new to our knowledge, in the sense that it is directly and exclusively based on the visibilities (L1a data), and thus can also be considered as an alternative to other brightness temperature recovery methods.

Index Terms— SMOS, MIRAS, RFI, non-differentiable convex optimization, total variation minimization.

1. INTRODUCTION

Meteorological and climate predictions are very sensitive to variables such as surface soil moisture (SSM) and sea surface salinity (SSS). The MIRAS instrument (Microwave Imaging Radiometer by Aperture Synthesis), carried on the SMOS satellite, is capable of measuring the corresponding brightness temperature of both SSM and SSS indirectly, in the L-band microwave, using interferometry and sensing the so-called *visibility function* [1], defined as the complex cross-correlation between the two signals collected by each pair (A_k, A_l) of antennas:

$$V_{k,l} = \frac{1}{\sqrt{\Delta_k \Delta_l}} \iint_{\|\xi\| \leq 1} U_k(\xi) U_l^*(\xi) (T_b(\xi) - T_r) \tilde{r}_{kl}(t) \frac{e^{-i2\pi \mathbf{u}_{kl}^T \xi}}{\sqrt{1 - \|\xi\|^2}} d\xi \quad (1)$$

Here, $\mathbf{u}_{kl}$ is the frequency baseline associated to A_k and A_l ; U_k, U_l are the corresponding normalized voltage patterns and Δ_k, Δ_l are the solid angles of the corresponding antennas. The Cartesian coordinates $\xi = (\xi_1, \xi_2)$ are the spatial do-

main coordinates, restricted to the unit circle. T_r is the physical temperature of the receivers (assumed the same for all receivers); $\tilde{r}_{kl}$ is the Fringe-Wash function, a function of the spatial delay $t = \frac{\mathbf{u}_{kl}^T \xi}{f_0}$, where $f_0 = \frac{c}{\lambda_0}$ is the central frequency of observation. Note that the brightness temperature T_b is a 2D function restricted to the unit circle ($\|\xi\| \leq 1$).

The MIRAS instrument is composed of three arms on a Y-shaped configuration, each arm carrying an array of antennas. This configuration leads to a hexagonal sampling grid of the visibility function. Figure 1 shows the star shaped domain Ω obtained from the MIRAS configuration.

If we denote by $T = T_b - T_r$, the samples of T in the hexagonal grid could be obtained from the visibility samples by solving the linear system $\mathbf{G}T = V$, where matrix $\mathbf{G}$ represents the discrete linear operator given by (1). Of course, this inverse problem is ill-posed since $\mathbf{G}$ is not invertible (due to the lack of information beyond Ω). Hence, additional constraints must be added to the model. In [2], Anterrieu proposes to solve it as a constrained least square problem, imposing that T has no frequency components outside Ω . This problem can be formulated as the unconstrained minimisation of $\|V - \mathbf{G}\mathbf{F}^*\mathbf{Z}_\Omega \hat{T}\|_2^2$ on $\hat{T}$, where $\mathbf{F}^*$ denotes the hexagonal Inverse Fourier Transform, $\mathbf{Z}_\Omega$ the zero padding operator and $\hat{T}$ the Fourier coefficients of T for frequencies in Ω . Let $\mathbf{J} = \mathbf{G}\mathbf{F}^*\mathbf{Z}_\Omega$, then $\hat{T} = \mathbf{J}^+V$ where $\mathbf{J}^+$ is the pseudo-inverse of $\mathbf{J}$: $\mathbf{J}^+ = (\mathbf{J}^*\mathbf{J})^{-1}\mathbf{J}^*$. This is the way the L1a product is obtained, and corresponds exactly to $\hat{T}$. In what follows we will note L1b data product as D_{L1b} .

Using L1b data, T can be recovered from D_{L1b} very easily: $T = \mathbf{F}^*\mathbf{Z}_\Omega D_{L1b}$. Naturally, as we have shown in [3], this simple Fourier inversion leads to potentially very strong Gibbs effects which are partially alleviated (as proposed by [2]) by the use of a Blackman window $\mathbf{B}$: $T = \mathbf{F}^*\mathbf{B}\mathbf{Z}_\Omega D_{L1b}$. As we pointed out in [3], such linear approaches are not well suited for the restoration of SMOS images, because usually the measurements are polluted by a number of strong outliers that correspond to signals emitted in the L-Band by illegal antennas¹. Because these outliers have frequencies beyond Ω , very strong Gibbs effects can be seen on the final brightness

¹International radio regulations reserve the L-band is exclusively to the Earth Exploration Satellite Service, space research and radio astronomy.

temperature images (see Figure 2). To eliminate them and to obtain a super-resolved solution, in [3] we proposed to solve the following variational formulation

$$\min_{u,o} \{TV(u) + \mu S(o)\} \text{ s.t. } \|W(F(o+u) - D_{L1b})\|_2^2 \leq |\Omega|\sigma^2, \quad (2)$$

where $TV(u)$ denotes the total variation semi-norm of the target image u , $S(o)$ a sparsity measure on the outliers (in practice ℓ_0 or ℓ_1 norms), W is a weighting matrix and σ^2 is the instrumental noise, assumed to be white and Gaussian (see [3] for details). Note that this method allows to automatically separate the restored image u from the outliers. Although this approach produced significantly better results than previous ones (Blackman and direct Fourier inversion), the model is not totally accurate since the restoration is performed starting from the D_{L1b} , which has already been subject to regularisation.

In this work, we improve our previous model, by directly dealing with visibilities (SMOS L1a data product). The goals are exactly the same than those of our previous work: to detect and remove signal effects generated from illegal emitters (outliers), while at the same time extrapolating the image spectrum in order to minimize Gibbs effects. As we will show in the following sections, this simple data change is not easy to implement because of the lack of regularization in the L1a data, and because of several issues that have to be considered to make the inverse problem numerically tractable. To our knowledge, this is the first work that tackles the problem directly and exclusively from the L1a product. Other approaches, such as [4], first localise the potential outliers based on the L1b product, then simulate the contribution on visibilities of each outlier to remove it from the original L1a data; finally they convert the outlier-free L1a data to L1b (by means of the pseudo-inverse matrix $\mathbf{J}^+$) as an intermediate step to recover the temperatures.

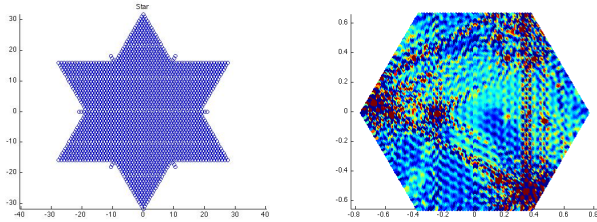


Fig. 1. Spectral domain Ω as- **Fig. 2.** Temperature obtained associated to MIRAS Y-shaped by the inverse Fourier transform of the L1b product.

2. PROPOSED APPROACH

As we already mentioned, visibilities and bright temperatures are related by the discrete linear operator given by (1). In

matrix notation², this is $\mathbf{G}u = V$. As before, the goal is to obtain the temperatures image u from the given visibilities V , knowing the non-invertibility of $\mathbf{G}$.

2.1. Variational formulation

We propose to recover u by solving the following constrained optimization problem:

$$\min_{u,o} \{TV(u) + \mu S(o)\} \text{ s.t. } \|\mathbf{G}(o+u) - V\|_2^2 \leq |\Omega|\sigma^2, \quad (3)$$

that can be reformulated as an unconstrained problem (see [3] for the details):

$$\min_{u,o} \{ \|\mathbf{G}(o+u) - V\|_2^2 + \lambda(TV(u) + \mu S(o)) \}. \quad (4)$$

The first term is the data fidelity term, which acts directly on the visibilities (L1a product). The second term acts as a regularizer and is the same as the one in our previous model; it is designed to separate the outliers o from the brightness temperature map u . $TV(u)$ intends to super-resolve u beyond spectral support Ω while avoiding Gibbs oscillations. The second term seeks to promote a sparse solution to the image of outliers by applying some sparsity operator to o . The parameter μ controls the trade-off between both terms; its choice can be formally derived from geometric considerations on the outliers. For the sake of brevity, we refer the reader to [3].

In order to reduce the “staircasing” effect inherent to many TV minimization methods, we make use of the *Spectral TV* introduced by Moisan [5] (see [3] for the implementation details). The final method can be stated as follows:

$$\min_{u,o} \left\{ \underbrace{\frac{1}{2} \|\mathbf{G}(o+u) - V\|_2^2}_{E_1(u,o)} + \underbrace{\lambda(TV_{\mathcal{H}}(u) + \mu S(o))}_{E_2(u,o)} \right\}, \quad (5)$$

where $TV_{\mathcal{H}}(u)$ denotes the *Spectral TV*.

2.2. Numerical implementation

Numerical minimization is based on a *Forward-Backward splitting* algorithm [6]. The k -th iteration starting from seed $x^0 = (u^0, o^0)$ is

$$\begin{cases} x^{k+1/2} &= x^k - \gamma \nabla E_1(x^k) \\ x^{k+1} &= \text{prox}_{\gamma E_2}(x^{k+1/2}). \end{cases}$$

In order to ensure convergence to the minimizer, γ must be smaller than $2/L$, where L is the Lipschitz constant of ∇E_1 (in our case $\gamma < 689$). Because the formulation only changes from the previous one on the data term E_1 , the proximal operator for E_2 remains the same as in the case of L1b data [3]:

$$\text{prox}_{\gamma E_2}(u, o) = \left(\text{prox}_{\gamma \lambda TV_{\mathcal{H}}}(u), \text{prox}_{\gamma \lambda \mu \|\cdot\|_1}(o) \right).$$

²For the sake of simplicity, we use the same symbol to refer to an image and its vectored form. Disambiguation follows easily from the context.

Now we differentiate E_1 and we obtain:

$$\nabla E_1(u, o) = (\mathbf{G}^* \mathbf{G}(u + o) - V, \mathbf{G}^* \mathbf{G}(u + o) - V).$$

$\mathbf{G}^* \mathbf{G}$ is a huge full matrix (16384×16384 since $\mathbf{G}$ is a 4695×16384 matrix). Explicit multiplication by this matrix at each iteration of the algorithm is computationally intractable. However, a change of basis to the Fourier domain changes the gradient term to

$$\nabla E_1(u, o) = F^*((\mathbf{G}F^*)^* \mathbf{G}F^* F(u + o) - (\mathbf{G}F^*)^* V)$$

which reveals an even larger (32768×32768) but highly sparse matrix $\mathbf{F} \mathbf{G}^* \mathbf{G} F^*$: to keep the energy at 99.99%, we only need to keep 0.0008% of the coefficients.

Finally, each iteration for the first step ($S(o) = \|o\|_1$) can be expressed as follows:

$$\begin{cases} u^{k+1/2} &= u^k - \gamma F^*(\mathbf{F} \mathbf{G}^* \mathbf{G} F^* F(u + o) - \mathbf{F} \mathbf{G}^* V) \\ o^{k+1/2} &= o^k - \gamma F^*(\mathbf{F} \mathbf{G}^* \mathbf{G} F^* F(u + o) - \mathbf{F} \mathbf{G}^* V) \\ u^{k+1} &= \text{prox}_{\gamma \lambda \text{TV}_{\mathcal{H}}}(u^{k+1/2}) \\ o^{k+1} &= s_{\gamma \lambda \mu}(o^{k+1/2}). \end{cases}$$

This iteration converges to a global minimum, that corresponds to the solution of problem (5) with sparsity operator $S(o) = \|o\|_1$. We use this solution as an initialisation for the second step, where the sparsity operator is non-convex, namely $S(o) = \|o\|_0$. For this problem, the same Forward-Backward method can be considered and is guaranteed to converge to a local minimizer [7]. Now, instead of the soft thresholding, the proximal operator for $S(o)$, which is now $\|o\|_0$, becomes the hard thresholding $h_{\sqrt{2\gamma\lambda\mu}}$.

3. EXPERIMENTS

To highlight the benefits of the proposed framework, we present two kinds of experiments. In the first one, we compare results from our approach to those obtained by previous works: a Fourier inversion of the L1b data, a simple Blackman apodization to smooth the outliers effects, and our previous L1b minimisation framework presented on [3]. Experiments are run on several snapshots from the SMOS dataset of march 2010. We have set σ equal to $5K$, which is the measurement error reported by the SMOS mission. Figures 4 and 5 show the results obtained for snapshots 996 and 1050. For geographic reference, figure 3 shows how these regions look in Google Earth from approximately the same viewing angles as the SMOS acquisitions.

Note that the acquired images are corrupted by several outliers that considerably degrade the data. It is clear that the present method outperforms both the direct inverse Fourier transform and the Blackman apodization. It also improves the results from our previous work based on L1b data, which resulted in more regularised images and thus, valuable details where lost. It is worth mentioning that the comparison of the results was confirmed by environmental scientists.



Fig. 3. Google Earth view of two of the regions used on the experiments. The left one corresponds to snapshot 996 and the right one to snapshot 1050.

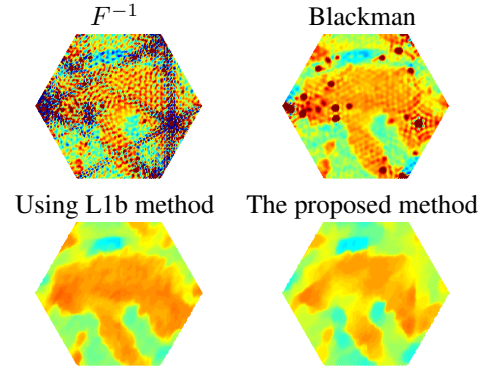


Fig. 4. Comparison between previous works and our method. This snapshot corresponds to Central Europe, with Italy clearly visible, and was taken on march 2010. Color scale was mapped between 0 and 300 Kelvin..

In the second experiment we perform a quantitative evaluation of our method, for which we need a *ground truth* data to compare with. We generated a simulated set in the following way. We first applied our algorithm to a real acquired image free of outliers. This leads to a denoised image u_{gt} that we selected as ground truth, and then we obtained from it the corresponding visibilities $V_{gt} = \mathbf{G}u_{gt}$. We then generated the outliers image by randomly selecting several pixels, and randomly generating intensity values between 300 and 30000 K. Similarly, this image T_δ was transformed into visibilities: $V_\delta = \mathbf{G}T_\delta$. Finally, we added noise that follows the given normal noise model with mean 0 and variance σ_v . The final simulated data was then $V_f = V_{gt} + V_\delta + n_v$, which follows exactly our image generation model. Figure 6 shows the results obtained on the simulated dataset. The table shows the estimation error on u for different error norms using both our previous method on L1b data, and the current one based on L1a. These errors were computed over the whole Fourier domain, since in the case of simulated images, no aliasing occurs. In all cases, the standard deviation of the residual measured in the spatial domain was 5 ± 0.2 , which is in agreement with the expected temperature noise model ($\sigma_T = 5.0$). We can see from this table a consistent improvement on the use

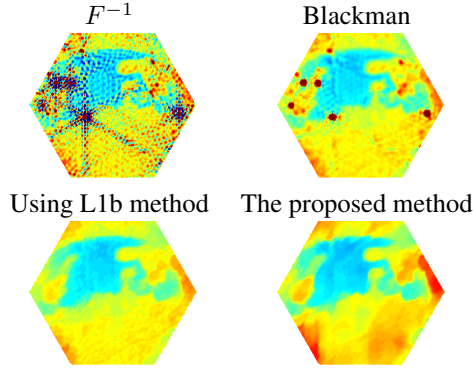


Fig. 5. Comparison between previous works and our method. This snapshot corresponds to North Europe and was taken on march 2010. Color scale were mapped between 0 and 300 Kelvin.

of the proposed method compared with the previous one.

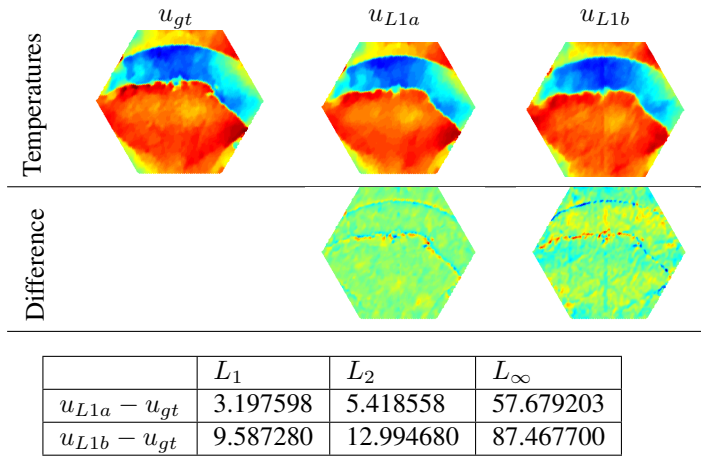


Fig. 6. Results from simulated data. The error is measured over all the image, not only the free of aliasing (FOA) zone.

4. CONCLUSIONS

In this article we present a new method to restore SMOS images directly and exclusively from the L1a product. Although the main objective of the method is to detect and remove the outliers produced by RFI, it can be considered as a new approach to obtain brightness temperature from the visibilities, even in the absence of outliers. The method follows the same methodology of our previous work [3], but adapted to the use of L1a instead of L1b data. We kept the previous idea of using two variables u and o to model brightness temperatures and outliers images separately, but the data term was adapted to use visibilities directly. This modification requires several changes: working directly with visibilities is hard since matrix $\mathbf{G}$ is numerically difficult to invert. We propose then to

work on the Fourier domain, which turns the problem more stable with a new matrix $\mathbf{GF}^*$ that becomes extremely sparse. This new approach is much closer to the real SMOS image formation model. This is not only a better model from a theoretical point of view: it gives better results because outliers are removed before any regularisation is done (in contrast with L1b methods where the RFI contributions are distributed all along the data). Consequently, the obtained images are less regularised when compared to L1b based methods, and small variations are kept, which is important since each pixel has a resolution between 30 and 50 km.

5. REFERENCES

- [1] I. Corbella, N. Duffo, M. Vall-llossera, A. Camps, and F. Torres, "The visibility function in interferometric aperture synthesis radiometry," *IEEE Trans. Geosci. Remote Sensing*, vol. 42, no. 8, pp. 1677–1682, 2004.
- [2] E. Anterrieu, "A resolving matrix approach for synthetic aperture imaging radiometers," *IEEE Transactions on Geoscience and Remote Sensing*, vol. 42, no. 8, pp. 1649 – 1656, aug. 2004.
- [3] J. Preciozzi, P. Muse, A. Almansa, S. Durand, F. Cabot, Y. Kerr, A. Khazaal, and B. Rouge, "Sparsity-based restoration of smos images in the presence of outliers," in *IEEE International Geoscience and Remote Sensing Symposium 2012*, 2012, pp. 3501–3504.
- [4] Rita Castro, Antonio Gutiérrez, and José Barbosa, "Iterative thresholding for sparse approximations," *IEEE Transactions on geoscience and remote sensing*, vol. 50, no. 5, pp. 1440–1447, May 2012.
- [5] L. Moisan, "How to discretize the total variation of an image?," *ICIAM07 Minisymposia*, 2007.
- [6] P. L. Combettes and V. R. Wajs, "Signal recovery by proximal forward-backward splitting," *SIAM Journal on Multiscale Modeling and Simulation*, vol. 4, no. 4, pp. 1168–1200, 2005.
- [7] Thomas Blumensath and Mike E Davies, "Iterative thresholding for sparse approximations," *The Journal of Fourier Analysis and Applications*, vol. 62, pp. 291–294, January 2008.