

Article

Bayesian 3D X-Ray Computed Tomography with a Hierarchical Prior Model for Sparsity in Haar Transform Domain

Li Wang *, Ali Mohammad-Djafari, Nicolas Gac and Mircea Dumitru

Laboratoire des signaux et système, Centralesupelec, CNRS, 3 Rue Joliot Curie, 91192 Gif sur Yvette, France; isct@free.fr (A.M-D.), nicolas.gac@l2s.centralesupelec.fr (N.G.); dumitru.i.mircea@gmail.com (M.D.)

* Correspondence: li.wang@l2s.centralesupelec.fr

Received: 26 October 2018; Accepted: 11 December 2018; Published: 16 December 2018



Abstract: In this paper, a hierarchical prior model based on the Haar transformation and an appropriate Bayesian computational method for X-ray CT reconstruction are presented. Given the piece-wise continuous property of the object, a multilevel Haar transformation is used to associate a sparse representation for the object. The sparse structure is enforced via a generalized Student- t distribution ($\mathcal{St}_g$), expressed as the marginal of a normal-inverse Gamma distribution. The proposed model and corresponding algorithm are designed to adapt to specific 3D data sizes and to be used in both medical and industrial Non-Destructive Testing (NDT) applications. In the proposed Bayesian method, a hierarchical structured prior model is proposed, and the parameters are iteratively estimated. The initialization of the iterative algorithm uses the parameters of the prior distributions. A novel strategy for the initialization is presented and proven experimentally. We compare the proposed method with two state-of-the-art approaches, showing that our method has better reconstruction performance when fewer projections are considered and when projections are acquired from limited angles.

Keywords: X-ray computed tomography; inverse problem; sparsity; hierarchical structure; generalized Student- t distribution; Haar transformation

1. Introduction

Computed Tomography (CT) has been developed and widely used in medical diagnosis [1] and industrial Non-Destructive Testing (NDT) [2] in recent decades. In CT, objects are observed using different techniques, for example X-rays [3], ultrasound [4], microwaves [5], or infra-red [6]. X-ray CT employs the absorption of X-rays by the organs in a body or by the materials in industrial components to reconstruct the internal structure of the imaged object. When performing X-ray CT, a set of X-ray images of the measured parts is acquired. The intensity measured by the X-ray images corresponds to the intensity of the radiation passing through and attenuated by the object. CT reconstruction is typically treated as an inverse problem.

The conventional analytical techniques for CT reconstruction are based on the Radon transform [7]. Filtered Back-Projection (FBP) [8] is the most frequently-used analytical method in practical applications. FBP performs well when reconstructing from sufficient data with a high signal-to-noise ratio (SNR), but it suffers from artifacts when reconstructing from insufficient data or with noise.

Owing to considerations regarding patients' health in medical CT and in order to reduce acquisition time in industrial applications, reconstruction with insufficient datasets is increasingly attracting the attention of researchers. Reconstruction from fewer projections is an ill-posed inverse problem [9,10]. In this case, conventional analytical reconstruction methods provide unsatisfactory

results, and iterative methods can be used to improve the reconstruction performance. The Algebraic Reconstruction Technique (ART) [11,12], the Simultaneous Algebraic Reconstruction Technique (SART) [13], and the Simultaneous Iterative Reconstruction Technique (SIRT) [14,15] are some of the iterative methods proposed initially. These methods consider the discretized forward system model: $\mathbf{g} = \mathbf{H}\mathbf{f}$, where $\mathbf{f} \in \mathbb{R}^{N \times 1}$ represents the object, $\mathbf{g} \in \mathbb{R}^{M \times 1}$ represents the observed dataset, and matrix $\mathbf{H} \in \mathbb{R}^{M \times N}$ is the linear projection operator, mainly based on the geometry of acquisition (e.g., parallel beam, cone beam, etc) [16–18]. Typically, the system of equations is under-determined, i.e., $N > M$. In this context, regularization methods are frequently used, and the forward system is modeled as:

$$\mathbf{g} = \mathbf{H}\mathbf{f} + \epsilon, \quad (1)$$

where $\epsilon \in \mathbb{R}^{M \times 1}$ represents the additive noise applied to the projection system. The regularization methods estimate the unknowns by minimizing a penalty criterion, which generally consists of two terms:

$$J(\mathbf{f}) = Q(\mathbf{g}, \mathbf{f}) + \lambda R(\mathbf{f}). \quad (2)$$

The loss function $Q(\mathbf{g}, \mathbf{f})$ describes discrepancies in the observed data, such as the quadratic (L2) loss $Q(\mathbf{g}, \mathbf{f}) = \|\mathbf{g} - \mathbf{H}\mathbf{f}\|_2^2$ or L_q loss $Q(\mathbf{g}, \mathbf{f}) = \|\mathbf{g} - \mathbf{H}\mathbf{f}\|_q^q$ with $1 \leq q < 2$. Other expressions such as the Huber function are also reported. The regularization term $R(\mathbf{f})$ is a penalty on the complement criterion of $\mathbf{f}$, such as restriction for smoothness $\|\Phi(\mathbf{f})\|_2^2$ or for sparsity $\|\Phi(\mathbf{f})\|_1$, where $\Phi(\mathbf{f})$ represents a linear function of $\mathbf{f}$. The parameter λ is known as the regularization parameter, which controls the trade-off between the forward model discrepancy and the penalty term.

By choosing different regularization functions $R(\mathbf{f})$, different regularization methods can be implemented. $R(\mathbf{f}) = 0$ refers to the Least-Squares (LS) method [19], with the drawback that the reconstruction is sensitive to the noise due to the ill-posedness of the problem and the ill-conditioning of the operator $\mathbf{H}$. Quadratic Regularization (QR), also known as the Tikhonov method [20], is given by $R(\mathbf{f}) = \|\Phi(\mathbf{f})\|_2^2$, where the linear operator $\Phi(\cdot)$ is the derivation operator in most cases. The well-known Total Variation (TV) method [21–23] is defined by $R(\mathbf{f}) = \|\mathbf{D}_{TV}\mathbf{f}\|_{TV}$ where $\mathbf{D}_{TV}$ is the gradient operator. $\mathbf{D}_{TV}$ is equal to $\|\mathbf{D}_x\mathbf{f}\|_1 + \|\mathbf{D}_y\mathbf{f}\|_1 + \|\mathbf{D}_z\mathbf{f}\|_1$ for a 3D object in an anisotropic form, where $\mathbf{D}_x$, $\mathbf{D}_y$ and $\mathbf{D}_z$ are respectively the gradient operators in the x , y and z directions. The L_1 norm is used in TV for sparse estimations, which enforces the sparsity of $\mathbf{D}_{TV}\mathbf{f}$. The appearance of the non-differentiable L_1 term leads to difficulties for the implementation of optimization algorithms. Many optimization algorithms have been proposed to solve this L_1 norm optimization problem, for example the primal-dual method [24], the split Bregman method [22], etc. In regularization optimization, due to the large projection data size and the great number of voxels, the explicit expression of the solution cannot be used directly because of the impossibility of inverting the large size matrix such as $(\mathbf{H}^T\mathbf{H} + \lambda\mathbf{D}^T\mathbf{D})^{-1}$. Hence, optimization algorithms such as gradient descent or conjugate gradient are often used.

More general regularization methods have been developed based on the constrained and dual-variable regularization method:

$$J(\mathbf{f}, \mathbf{z}) = Q_1(\mathbf{g}, \mathbf{f}) + \eta Q_2(\mathbf{f}, \mathbf{z}) + \lambda R(\mathbf{z}), \quad (3)$$

which corresponds to the maximum a posterior optimization of a hierarchical structured model where both $\mathbf{f}$ and $\mathbf{z}$ are unknown variables. In such a model, the penalty regularization term is set on $\mathbf{z}$, which is associated with $\mathbf{f}$ via a linear transformation. The loss functions $Q_1(\mathbf{g}, \mathbf{f})$ and $Q_2(\mathbf{f}, \mathbf{z})$ are for example quadratic (L2), i.e., $Q_1(\mathbf{g}, \mathbf{f}) = \|\mathbf{g} - \mathbf{H}\mathbf{f}\|_2^2$ and $Q_2(\mathbf{f}, \mathbf{z}) = \|\mathbf{f} - \mathbf{D}\mathbf{z}\|_2^2$ where $\mathbf{D}$ is a linear transform operator such as a wavelet transformation.

Among the methods treating this type of regularization problem, we mention here the Alternating Direction Method of Multipliers (ADMM) [25]. It minimizes $\Phi(\mathbf{f}) + \Psi(\mathbf{z})$ subject to $\mathbf{A}\mathbf{f} + \mathbf{B}\mathbf{z} = \mathbf{C}$, and it covers a large number of estimation forms. One example is when $\Phi(\mathbf{f}) = \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2$, $\Psi(\mathbf{z}) = R(\mathbf{z})$,

$A = I$, $B = -D$, and $C = 0$ and refers to the above-mentioned bi-variable regularization method corresponding to Equation (3).

In the above-mentioned regularization methods, there is always a regularization parameter λ to be fixed. Sometimes, the regularization term consists of more than one parts, and each of them are weighted by a regularization parameter, for example the elastic-net regularizer [26]. In these cases, two or even more regularization parameters need to be fixed. Cross Validation (CV) and the L-curve method [20,27,28] are conventional methods used to determine suitable values for these parameters. However, this work must be repeated for different simulated datasets and is therefore very costly. Statistical methods, therefore, have been developed and used to solve this problem.

From the probabilistic point of view, a Gaussian model for the additive noise in the forward model, Equation (1), leads to the quadratic expression $\|g - Hf\|^2$ in the corresponding regularization criterion. However, in some types of tomography, for example Positron Emission Tomography (PET) or X-ray tomography with a very low number of phantoms, the noise is modeled by a Poisson distribution. In order to account for a more precise modeling of the noise and the other variables and parameters, statistical methods are used [29]. The Maximum Likelihood (ML) methods [30] and different estimation algorithms such as the Expectation Maximization (EM) algorithms [31], the Stochastic EM (SEM) [32], or the Ordered Subsets-EM (OS-EM) [33] are commonly used in PET-CT reconstruction problems.

Another widely-used type of probabilistic method for PET or X-ray CT reconstruction is the Bayesian inference [34–37]. The prior knowledge is translated by the prior probability model and is used to obtain the expression of the posterior distribution. The basic Bayesian formula is:

$$p(f|g, \theta) = \frac{p(g|f, \theta_1)p(f|\theta_2)}{p(g|\theta)}, \text{ with } p(g|\theta) = \int p(g|f, \theta_1)p(f|\theta_2) df, \quad (4)$$

where $p(g|f, \theta_1)$ is the likelihood, $p(f|\theta_2)$ is the prior distribution, $p(f|g, \theta)$ is the posterior distribution, $\theta = (\theta_1, \theta_2)$ are the parameters of these different distributions, and $p(g|\theta)$ is the evidence of the parameters in the data g . By using Maximum A Posterior (MAP) estimator $\hat{f} = \arg \max_f \{p(f|g, \theta)\} = \arg \min_f \{-\ln p(f|g, \theta)\}$, links between the Bayesian method and almost all the regularization methods can be illustrated. A Gaussian prior for $p(f)$ in Equation (4) leads to the quadratic (L_2) regularization method, while a Laplacian prior in Equation (4) leads to the L_1 (LASSO or TV) regularization method. The regularization parameter can be related to θ_1 and θ_2 . One advantage of the Bayesian method is having some explanation for the regularization parameter via its link with θ_1 and θ_2 . For example, when $p(g|f, \theta_1)$ and $p(f|\theta_2)$ are Gaussian with θ_1 and θ_2 , respectively the variances of the noise and the variance of the prior, then the regularization parameter is $\lambda = \theta_1/\theta_2$. Another advantage of the Bayesian method is that these parameters can also be estimated to achieve unsupervised or semi-supervised methods. This is achieved by obtaining the expression of the joint posterior probability law:

$$p(f, \theta|g) = \frac{p(g|f, \theta_1)p(f|\theta_2)p(\theta)}{p(g)}, \quad (5)$$

where $p(\theta)$ is an appropriate prior on θ . For a hierarchical structured model where a hidden variable z appears in the prior model, we have:

$$p(f, z, \theta|g) = \frac{p(g|f, \theta_1)p(f|z, \theta_2)p(z|\theta_3)p(\theta)}{p(g)}, \quad (6)$$

where $\theta = [\theta_1, \theta_2, \theta_3]$.

With the posterior distribution obtained from an unsupervised Bayesian inference as in Equation (5), we distinguish three estimation methods. The first method consists of integrating out θ from $p(f, \theta|g)$ to obtain $p(f|g)$ and then using $p(f|g)$ to infer on f . The second approach is firstly to marginalize $p(f, \theta|g)$ with respect to f to obtain $p(\theta|g) = \int p(f, \theta|g) df$ and estimate $\hat{\theta} = \arg \max_{\theta} \{p(\theta|g)\}$, then use $\hat{\theta}$ as it was known. Unfortunately, these approaches do not often

give explicit expressions for $p(f|g)$ or $p(\theta|g)$. The third and easiest algorithm to implement is the joint optimization, which estimates variable f and parameter θ iteratively and alternately. Bayesian point estimators such as Joint Maximum A Posteriori (JMAP) [38] and Posterior Mean (PM) [39] via Variational Bayesian Approximation (VBA) methods [40–42] are often used.

In order to distinguish the details of a reconstructed object, a high-resolution image is expected. In industrial applications, especially for the NDT of a large-size object, the size of the projection (1000 images of 1000^2 pixels) and the number of voxels (1000^3 voxels) become critical, and so does the projection and back-projection operators in CT. It is necessary to account for some computational aspects, for example the GPU processor [43,44].

In our previous work [45], we proposed to use the Bayesian method via a synthesis model, in which the multilevel Haar transformation coefficient z of the image is first estimated, and then, the final image reconstruction result is obtained from post-processing: $\hat{f} = D\hat{z}$. In this case, when using a Laplacian prior model and the MAP estimator, the problem becomes equivalent to the optimization of $J(z) = \|g - HDz\|^2 + \lambda \|z\|_1$, which is a typical L_1 regularization method. The particularity of our work was to use a generalized Student- t (St_g) prior model [46] in place of the Laplacian model.

In this paper, we present a Hierarchical Haar transform-based Bayesian Method (HHBM), first proposed in [47], in which the object to be reconstructed, f , is related to the Haar transformation coefficient z by $f = Dz + \xi$ where ξ represents the modelization uncertainties. f and z are estimated simultaneously. Wavelets provide an optimal representation for a piecewise continuous function consisting of homogeneous blocs separated by jump discontinuities (the contours), as the wavelet representation is sparse for such signals. The transformations used are, for example, the Haar transformation [48], the Curvelet Transformation (CVT) [49], the Contourlet Transformation (CT) [50], Dual-Tree Complex Wavelet Transform (DT-CWT) [51], etc. As long as the object under consideration, f , is piecewise continuous or constant, the Haar transform is appropriate, with the advantage that: (1) the transform coefficients are sparse; (2) the transformation operator is orthogonal so that the inverse operator and the transpose are identical; and (3) the computation of this transformation consists of only additions and subtractions, while the cost of computation is only $O(\sqrt{N})$ where N is the size of the object f .

The sparsity of the transformation coefficient is generally defined by three classes of distributions: the generalized Gaussian distributions [52], the mixture distributions [53], and the heavy-tailed distributions [54]. In this paper, we use a generalization of the Student- t distribution (St_g), which belongs to the heavy-tailed family and has many advantages when enforcing the sparsity of variables [46].

In this paper, we extend extensively the previous work by: (1) adapting the forward model and prior models to the 3D case, which is more appropriate for real 3D large data size applications; (2) comparing the RMSE of the phantom reconstructed by the HHBM method with those by the conventional QR and TV methods, we show the advantages of the semi-supervised property of the HHBM method and that the HHBM method outperforms the TV method when insufficient data are estimated; (3) proposing new ideas for fixing the hyper-parameters in the proposed model; and (4) evaluating the performance of the proposed method in the situations when the number or the angle distribution of the projections is limited.

The rest of this paper is organized as follows: Section 2 presents the proposed hierarchically-structured Bayesian method; Section 3 gives the details of the implementations and the choice of hyperparameters, as well as the simulation results; some points on the initialization of hyper-parameters are discussed in Section 4. Conclusions are drawn and prospective future research is presented in Section 5.

2. The Semi-Supervised Hierarchical Model

The Hierarchical Haar-based Bayesian Method (HHBM) is presented, in which the object f and its multilevel Haar transform coefficient z are considered jointly. A sparse enforcing prior is defined

on $\mathbf{z}$. The wavelet transformation has been used for the reconstruction of tomography images in some state-of-the-art works [55–58], using both regularization and Bayesian methods. In these state-of-the-art methods, the phantom $\mathbf{f}$ is obtained by a post-processing from reconstructed coefficient $\mathbf{z}$. In this paper, the phantom $\mathbf{f}$ and the coefficient $\mathbf{z}$ are simultaneously estimated.

2.1. Forward System Model and Likelihood

In the proposed method, the forward model introduced in Equation (1) is considered. Generally, the noise in tomography is modeled by a Poisson distribution [59], but in X-ray CT, the Gaussian approximation is often used. We adopt the Gaussian approximation and propose to use a zero mean and non-stationary model where the variance is considered to be unknown, belonging to an inverse Gamma distribution given that this distribution provides a good adaptation of the positivity property of the variances v_{ϵ_i} :

$$p(\epsilon|\mathbf{v}_\epsilon) = \mathcal{N}(\epsilon|\mathbf{0}, \mathbf{V}_\epsilon), \quad \mathbf{V}_\epsilon = \text{diag}[\mathbf{v}_\epsilon], \quad \text{where } \mathbf{v}_\epsilon = [v_{\epsilon_1}, \dots, v_{\epsilon_M}]' \in \mathbb{R}^{M \times 1} \quad (7)$$

$$p(\mathbf{v}_\epsilon|\alpha_{\epsilon_0}, \beta_{\epsilon_0}) = \prod_{i=1}^M \mathcal{IG}(v_{\epsilon_i}|\alpha_{\epsilon_0}, \beta_{\epsilon_0}) \quad (8)$$

The vector $\mathbf{v}_\epsilon$ is considered in order to account for the difference of sensitivity to noise for each detector in each projection direction.

According to the forward model of the linear system, Equation (1), and the prior model of the noise, Equation (7), the likelihood of this model system is:

$$p(\mathbf{g}|\mathbf{f}, \mathbf{v}_\epsilon) = \mathcal{N}(\mathbf{g}|\mathbf{H}\mathbf{f}, \mathbf{V}_\epsilon). \quad (9)$$

In Bayesian inference, the likelihood is combined with the prior distributions to determine the posterior distribution.

2.2. Hierarchical Prior Model and Prior Distributions

Typically, the objects considered in medical and industrial X-ray CT are piecewise continuous. In this paper, a hierarchical prior model is used to define the piecewise continuous property. In this hierarchical prior model, a sparsity enforcing model is defined for the wavelet transformation coefficients of the image. A large number of methods accounting for the sparse structure of the solution have been proposed in the literature. Among them, the L1 regularization method is most frequently used, which minimizes the criterion $J(\mathbf{f}) = \|\mathbf{g} - \mathbf{H}\mathbf{f}\|_2^2 + \lambda \|\mathbf{\Phi}\mathbf{f}\|_1$ where $\mathbf{\Phi}$ is a linear operator, for example the gradient in the TV method. Another class of methods, known as “synthesis” [60], minimizes $J(\mathbf{z}) = \|\mathbf{g} - \mathbf{H}\mathbf{D}\mathbf{z}\|_2^2 + \lambda \|\mathbf{z}\|_1$ where $\mathbf{f} = \mathbf{D}\mathbf{z}$, and $\mathbf{z} = \mathbf{D}^{-1}\mathbf{f}$ is for example a wavelet transform.

In this paper, we propose to use the multilevel Haar transformation as the sparse dictionary. The transformation is modeled using a discretized forward model:

$$\mathbf{f} = \mathbf{D}\mathbf{z} + \boldsymbol{\xi}, \quad (10)$$

where $\mathbf{D} \in \mathbb{R}^{N \times N}$ represents the inverse multilevel Haar transformation and $\boldsymbol{\xi} \in \mathbb{R}^{N \times 1}$ represents the uncertainties of the transformation, which is introduced to relax the exact relation of the transform operator $\mathbf{D}$. $\boldsymbol{\xi}$ is supposed to be sparse. Unlike the gradient operator used in the TV method, the multilevel Haar transformation is orthogonal, i.e. $\mathbf{D}^{-1} = \mathbf{D}^T$. This property provides certain advantages during optimization, especially for the big data size problems, as the inversion and transpose of the operator are identical and can be replaced by each other for different types of $\mathbf{D}$.

$\boldsymbol{\xi}$ is modeled by a Gaussian distribution, $p(\boldsymbol{\xi}) = \mathcal{N}(\boldsymbol{\xi}|\mathbf{0}, \mathbf{V}_\xi)$, where $\mathbf{V}_\xi = \text{diag}[\mathbf{v}_\xi]$ and $\mathbf{v}_\xi = [v_{\xi_1}, \dots, v_{\xi_N}]'$. $\mathbf{v}_\xi$ is considered an unknown variance. It is modeled in order to realize a semi-supervised system where the variance is estimated. Here, v_ξ is modeled by an inverse Gamma

distribution, with the same consideration as the model of v_ϵ . The Gaussian model with an inverse Gamma distributed variance, $p(v_\xi|\alpha_{\xi_0}, \beta_{\xi_0}) = \prod_{j=1}^N \mathcal{IG}(v_{\xi_j}|\alpha_{\xi_0}, \beta_{\xi_0})$, leads to a generalized Student- t (St_g) distribution [46]. Consequently, a St_g distribution is derived for ξ , and the sparse property of ξ can be guaranteed. From Equation (10), the conditional distribution $p(f|z, v_\xi)$ is derived:

$$p(f|z, v_\xi) = \mathcal{N}(f|Dz, V_\xi), \quad (11)$$

with:

$$p(v_\xi|\alpha_{\xi_0}, \beta_{\xi_0}) = \prod_j^N \mathcal{IG}(v_{\xi_j}|\alpha_{\xi_0}, \beta_{\xi_0}). \quad (12)$$

For practical applications where these parameters are not known or difficult to obtain, we use a semi-supervised method in which the variances of noises, v_ϵ and v_ξ , are unknown. In HHBM, the inverse Gamma distribution is used to model v_ϵ and v_ξ , $p(v_\epsilon) = \mathcal{IG}(v_\epsilon|\alpha_{\epsilon_0}, \beta_{\epsilon_0})$ and $p(v_\xi) = \mathcal{IG}(v_\xi|\alpha_{\xi_0}, \beta_{\xi_0})$. Consequently, both ξ and ϵ are modeled by a St_g distribution.

Vector $z = [z_1, \dots, z_N]'$ represents the multilevel Haar transformation coefficient of piece-wise continuous f . As mentioned above, z is sparse. In this paper, the generalized Student- t distribution (St_g) [46] is used to enforce the sparsity structure of z . The St_g distribution can be expressed as the marginal of a bivariate normal-inverse Gamma distribution:

$$St_g(z|\alpha, \beta) = \int \mathcal{N}(z|0, v) \mathcal{IG}(v|\alpha, \beta) dv. \quad (13)$$

Thanks to the fact that normal and inverse Gamma are conjugate distributions, the use of the St_g via Equation (13) simplifies the computations when using the Bayesian point estimators such as the posterior mean via the Variational Bayesian Approximation (VBA) method [42].

From Equation (13), the St_g prior distribution modeling z is expressed as the following model:

$$p(z|v_z) = \mathcal{N}(z|0, V_z), \text{ where } V_z = \text{diag}[v_z], \quad v_z = [v_{z_1}, v_{z_2}, \dots, v_{z_N}] \quad (14)$$

$$p(v_z|\alpha_{z_0}, \beta_{z_0}) = \prod_j^N \mathcal{IG}(v_{z_j}|\alpha_{z_0}, \beta_{z_0}), \quad (15)$$

where $v_{z_j}, \forall j = 1 : N$ are i.i.d. distributed. The difference between the standard St distribution, $St(z|\nu) = \int \mathcal{N}(z|0, v) \mathcal{IG}(v|\frac{\nu}{2}, \frac{\nu}{2}) dv$, and the generalized St_g , given in Equation (13), is that $St(z|\nu)$ is governed by one parameter ν , but $St_g(z|\alpha, \beta)$ is governed by two parameters (α, β) . With these two parameters, the St_g does not only enforce the sparsity of the variable, but also controls the sparsity rate [46]. By changing the values of the two hyper-parameters α_{z_0} and β_{z_0} , we can obtain either a heavy-tailed distribution with a narrow peak or a distribution approaching a Gaussian distribution.

2.3. The HHBM Method

The prior models of the proposed Bayesian method based on the forward model of Equation (1) and the prior model of Equation (10) are:

$$p(\mathbf{g}|\mathbf{f}, \mathbf{v}_\epsilon) \propto |\mathbf{V}_\epsilon|^{-\frac{1}{2}} \exp \left[-\frac{1}{2} (\mathbf{g} - \mathbf{H}\mathbf{f})^T \mathbf{V}_\epsilon^{-1} (\mathbf{g} - \mathbf{H}\mathbf{f}) \right], \quad (16)$$

$$p(\mathbf{f}|\mathbf{z}, \mathbf{v}_\xi) \propto |\mathbf{V}_\xi|^{-\frac{1}{2}} \exp \left[-\frac{1}{2} (\mathbf{f} - \mathbf{D}\mathbf{z})^T \mathbf{V}_\xi^{-1} (\mathbf{f} - \mathbf{D}\mathbf{z}) \right], \quad (17)$$

$$p(\mathbf{z}|\mathbf{v}_z) \propto |\mathbf{V}_z|^{-\frac{1}{2}} \exp \left[-\frac{1}{2} \mathbf{z}^T \mathbf{V}_z^{-1} \mathbf{z} \right], \quad (18)$$

$$p(\mathbf{v}_z|\alpha_{z_0}, \beta_{z_0}) \propto \prod_j^N v_{z_j}^{-(\alpha_{z_0}+1)} \exp \left[-\beta_{z_0} v_{z_j}^{-1} \right], \quad (19)$$

$$p(\mathbf{v}_\epsilon|\alpha_{\epsilon_0}, \beta_{\epsilon_0}) \propto \prod_i^M v_{\epsilon_i}^{-(\alpha_{\epsilon_0}+1)} \exp \left[-\beta_{\epsilon_0} v_{\epsilon_i}^{-1} \right], \quad (20)$$

$$p(\mathbf{v}_\xi|\alpha_{\xi_0}, \beta_{\xi_0}) \propto \prod_j^N v_{\xi_j}^{-(\alpha_{\xi_0}+1)} \exp \left[-\beta_{\xi_0} v_{\xi_j}^{-1} \right]. \quad (21)$$

Figure 1 shows the generative graph of the proposed model in which the hyperparameters in the rectangles need to be initialized:

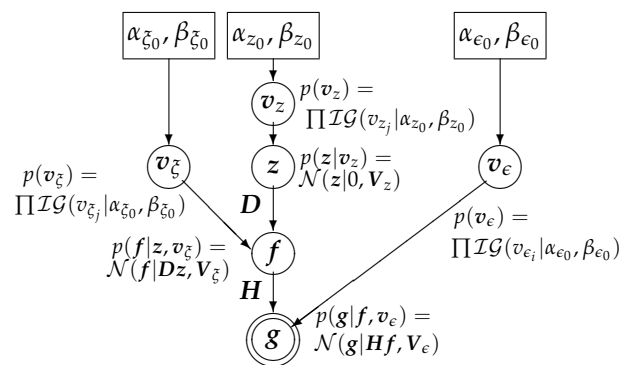


Figure 1. Generative graph of the proposed model illustrating all the unknowns (circles), hyperparameters (boxes), and data (double circles).

Via the Bayes rule, Equation (6), the joint posterior distribution of all the unknowns given in the data is derived:

$$\begin{aligned} p(\mathbf{f}, \mathbf{z}, \mathbf{v}_\epsilon, \mathbf{v}_\xi, \mathbf{v}_z | \mathbf{g}) &= \frac{p(\mathbf{g}, \mathbf{f}, \mathbf{z}, \mathbf{v}_\epsilon, \mathbf{v}_\xi, \mathbf{v}_z)}{p(\mathbf{g})} \\ &= \frac{p(\mathbf{g}|\mathbf{f}, \mathbf{v}_\epsilon) p(\mathbf{f}|\mathbf{z}, \mathbf{v}_\xi) p(\mathbf{z}|\mathbf{v}_z) p(\mathbf{v}_z) p(\mathbf{v}_\epsilon) p(\mathbf{v}_\xi)}{p(\mathbf{g})} \\ &\propto p(\mathbf{g}|\mathbf{f}, \mathbf{v}_\epsilon) p(\mathbf{f}|\mathbf{z}, \mathbf{v}_\xi) p(\mathbf{z}|\mathbf{v}_z) p(\mathbf{v}_z) p(\mathbf{v}_\epsilon) p(\mathbf{v}_\xi). \end{aligned} \quad (22)$$

Bayesian point estimators are often used for estimation via the a posteriori distribution. In this paper, we focus on the JMAP estimation, given that in the case of the large data size of the 3D object, the computational costs for the VBA algorithm is too expensive.

2.4. Joint Maximum a Posteriori Estimation

The negative logarithm of the posterior distribution is used as the criterion of optimization in order to simplify the exponential terms. The maximization of the posterior distribution becomes a minimization of the criterion:

$$\begin{aligned} (f, z, v_z, v_\epsilon, v_\xi) &= \arg \max \{p(f, z, v_\epsilon, v_\xi, v_z | g)\} \\ &= \arg \min \{-\ln p(f, z, v_\epsilon, v_\xi, v_z | g)\} \\ &= \arg \min J(f, z, v_z, v_\epsilon, v_\xi). \end{aligned} \quad (23)$$

We substitute the distribution formulas and obtain:

$$\begin{aligned} J(f, z, v_z, v_\epsilon, v_\xi) &= -\ln p(f, z, v_\epsilon, v_\xi, v_z | g) \\ &= \frac{1}{2} \sum_i^M \ln v_{\epsilon_i} + \frac{1}{2} (g - Hf)^T V_\epsilon^{-1} (g - Hf) \\ &\quad + \frac{1}{2} \sum_j^N \ln v_{\xi_j} + \frac{1}{2} (f - Dz)^T V_\xi^{-1} (f - Dz) \\ &\quad + \frac{1}{2} \sum_j^N \ln v_{z_j} + \frac{1}{2} z^T V_z^{-1} z + (\alpha_{z_0} + 1) \sum_j^N \ln v_{z_j} \\ &\quad + \beta_{z_0} \sum_j^N v_{z_j}^{-1} + (\alpha_{\epsilon_0} + 1) \sum_i^M \ln v_{\epsilon_i} + \beta_{\epsilon_0} \sum_i^M v_{\epsilon_i}^{-1} \\ &\quad + (\alpha_{\xi_0} + 1) \sum_j^N \ln v_{\xi_j} + \beta_{\xi_0} \sum_j^N v_{\xi_j}^{-1}. \end{aligned} \quad (24)$$

The unknown variables are determined by obtaining the expressions of the alternate minimization in Equation (24):

$$\hat{f} = \left(H^T \hat{V}_\epsilon^{-1} H + \hat{V}_\xi^{-1} \right)^{-1} \left(H^T \hat{V}_\epsilon^{-1} g + \hat{V}_\xi^{-1} D \hat{z} \right), \quad (25)$$

$$\hat{z} = \left(D^T \hat{V}_\xi^{-1} D + \hat{V}_z^{-1} \right)^{-1} D^T \hat{V}_\xi^{-1} \hat{f}, \quad (26)$$

$$\hat{v}_{z_j} = \left(\beta_{z_0} + \frac{1}{2} \hat{z}_j^2 \right) / (\alpha_{z_0} + 3/2), \quad (27)$$

$$\hat{v}_{\epsilon_i} = \left(\beta_{\epsilon_0} + \frac{1}{2} \left(g_i - [H\hat{f}]_i \right)^2 \right) / (\alpha_{\epsilon_0} + 3/2), \quad (28)$$

$$\hat{v}_{\xi_j} = \left(\beta_{\xi_0} + \frac{1}{2} \left(\hat{f}_j - [D\hat{z}]_j \right)^2 \right) / (\alpha_{\xi_0} + 3/2), \quad (29)$$

$\forall i \in [1, M]$ and $\forall j \in [1, N]$.

In 3D X-ray CT, the inversion of matrix $\left(H^T \hat{V}_\epsilon^{-1} H + \hat{V}_\xi^{-1} \right)^{-1}$ and $\left(D^T \hat{V}_\xi^{-1} D + \hat{V}_z^{-1} \right)^{-1}$ in Equations (25) and (26) is impossible due to the large data size. First-order optimization methods are generally used in this case. In this paper, we use the gradient descent algorithm:

$$\text{for } k = 1 \rightarrow I_G : \hat{f}^{(k+1)} = \hat{f}^{(k)} - \hat{\gamma}_f^{(k)} \nabla J_f(\hat{f}^{(k)}), \quad (30)$$

$$\text{for } k = 1 \rightarrow I_G : \hat{z}^{(k+1)} = \hat{z}^{(k)} - \hat{\gamma}_z^{(k)} \nabla J_z(\hat{z}^{(k)}), \quad (31)$$

where I_G is the number of iterations for the gradient descent estimation and $\nabla J_f(\cdot)$ and $\nabla J_z(\cdot)$ are the derivatives of the criterion (24) regarding f and z , respectively. $\hat{\gamma}_f(\cdot)$ and $\hat{\gamma}_z(\cdot)$ are the corresponding descent step lengths, which are obtained by using an optimized step length strategy [61]:

$$\nabla J(f) = -H^T V_\epsilon^{-1} (g - Hf) + V_\xi^{-1} (f - Dz), \quad (32)$$

$$\nabla J(z) = -D^T V_\xi^{-1} (f - Dz) + V_z^{-1} z, \quad (33)$$

$$\hat{\gamma}_f^{(k)} = \frac{\|\nabla J(\hat{f}^{(k)})\|^2}{\|\hat{Y}_\epsilon H \nabla J(\hat{f}^{(k)})\|^2 + \|\hat{Y}_\xi \nabla J(\hat{f}^{(k)})\|^2}, \quad (34)$$

$$\hat{\gamma}_z^{(k)} = \frac{\|\nabla J(\hat{z}^{(k)})\|^2}{\|\hat{Y}_\xi D \nabla J(\hat{z}^{(k)})\|^2 + \|\hat{Y}_z \nabla J(\hat{z}^{(k)})\|^2}, \quad (35)$$

where $Y_\epsilon = V_\epsilon^{-\frac{1}{2}}$, $Y_\xi = V_\xi^{-\frac{1}{2}}$, and $Y_z = V_z^{-\frac{1}{2}}$.

The algorithm concerning the optimization of all the unknowns is given in Algorithm 1.

Algorithm 1 The JMAP Algorithm for the HHBM Method

- 1: Fix parameters $\alpha_{z_0}, \beta_{z_0}, \alpha_{\epsilon_0}, \beta_{\epsilon_0}, \alpha_{\xi_0}, \beta_{\xi_0}, l$
 - 2: **Input:** H, D, g
 - 3: **Output:** $\hat{f}, \hat{z}, \hat{v}_z, \hat{v}_\epsilon, \hat{v}_\xi$
 - 4: **Initialization:**
 - 5: $\hat{f} \leftarrow$ normalized FBP
 - 6: $\hat{z} \leftarrow l$ -level Haar transformation of $\hat{f}$
 - 7: **for** $k = 1$ to I_{max} **do**
 - 8: $\hat{f}^{(0)} = \hat{f}$
 - 9: **for** $k = 1$ to I_G **do**
 - 10: Calculate $\nabla J(\hat{f}^{(k-1)})$ according to Equation (32)
 - 11: Update $\hat{\gamma}_f^{(k)}$ according to Equation (34)
 - 12: Update $\hat{f}^{(k)} = \hat{f}^{(k-1)} - \hat{\gamma}_f^{(k)} \nabla J(\hat{f}^{(k-1)})$
 - 13: **end for**
 - 14: $\hat{f} = \hat{f}^{(I_G)}$
 - 15: $\hat{z}^{(0)} = \hat{z}$
 - 16: **for** $k = 1$ to I_G **do**
 - 17: Calculate $\nabla J(\hat{z}^{(k-1)})$ according to Equation (33)
 - 18: Update $\hat{\gamma}_z^{(k)}$ according to Equation (35)
 - 19: Update $\hat{z}^{(k)} = \hat{z}^{(k-1)} - \hat{\gamma}_z^{(k)} \nabla J(\hat{z}^{(k-1)})$
 - 20: **end for**
 - 21: $\hat{z} = \hat{z}^{(I_G)}$
 - 22: Optimize v_z according to Equation (27)
 - 23: Optimize v_ϵ according to Equation (28)
 - 24: Optimize v_ξ according to Equation (29)
 - 25: **end for**
-

3. Initialization and Experimental Results

For the simulations, the 3D simulated “Shepp–Logan” phantom, shown in Figure 2 on the left, Figure 3 on the top, and the 3D real “head” object, shown in Figure 2 on the right, Figure 3 on the bottom, both of size 256^3 , are used as the objects of interest to compare the performance of the proposed method to the performance of the other state-of-the-art methods. Both the Shepp–Logan and head phantoms consist of several different homogeneous areas, so both are piecewise continuous. The voxel values of the original objects are normalized to $[0, 1]$. The projection directions are uniformly distributed, and each projection consists of 256^2 detectors corresponding to a 256^2 sized image.

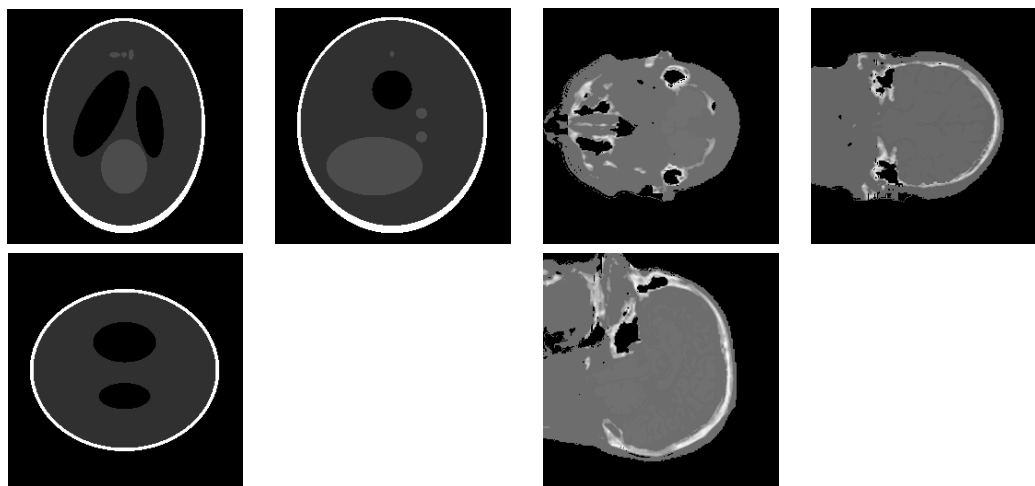


Figure 2. The three figure on the left show the three middle slice views of the three-dimensional Shepp–Logan phantom, and the three figures on the right show the three-dimensional head phantom.

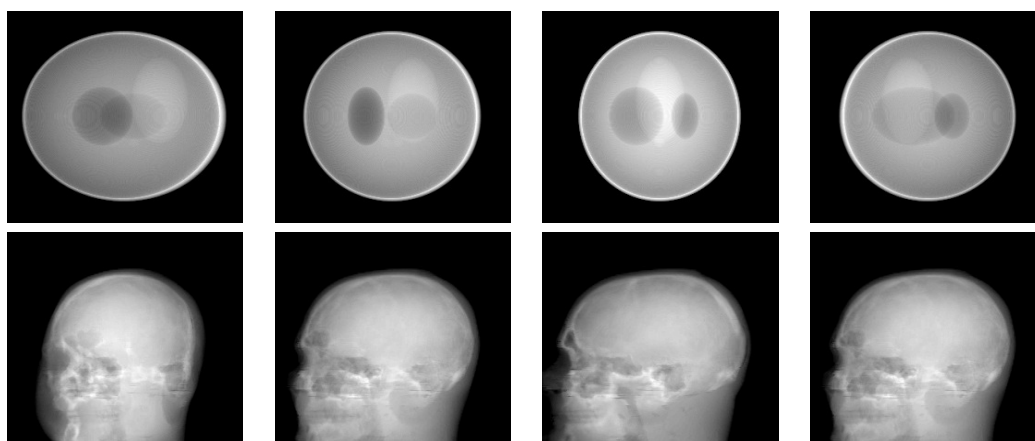


Figure 3. (top) Four projections of the 3D Shepp–Logan phantom from 30, 60, 90, 120 degrees (left to right, respectively); (bottom) four projections of the 3D head phantom from 30, 60, 90, 120 degrees (left to right, respectively).

The proposed HHBM method is compared with the conventional Quadratic Regularization (QR) and Total Variation (TV) methods. For the QR method, the gradient descent algorithm is used for the 3D large data size problem. For the TV method, the split Bregman method [22] is used to solve the L_1 norm minimization problem.

To evaluate the proposed method and compare it with the state-of-the-art methods, four different metrics are used:

- the Relative Mean Squared Error (RMSE), $RMSE = \|\mathbf{f} - \hat{\mathbf{f}}\|^2 / \|\mathbf{f}\|^2$, which shows a relative error of the results;
- the Improvement of the Signal-to-Noise Ratio (ISNR), which measures the improvement during iterations;
- the Peak Signal-to-Noise-Ratio (PSNR), which presents the SNR relative to the peak data value;
- the Structural Similarity of IMage (SSIM) [62], which evaluates the quality of the result approaching human vision.

In 3D X-ray CT, the projection matrix $\mathbf{H}$ is very large and is not accessible. For the simulations, only the projection operator $\mathbf{H}\mathbf{f}$ and the back-projection operator $\mathbf{H}^T\mathbf{g}$ are used. Considering that the costly projection and back-projection operators are computed in every iteration, the GPU processor is used via the ASTRA toolbox [63] to accelerate the computation.

3.1. Initializations

The initialization for the variables $\mathbf{f}$ and $\mathbf{z}$, as well as the hyperparameters α_{ϵ_0} , β_{ϵ_0} , α_{ξ_0} , β_{ξ_0} , α_{z_0} , and β_{z_0} are discussed in this section.

The reconstructed phantom obtained by using the Filtered Back Projection (FBP) method is considered to be initial value $\hat{\mathbf{f}}_{ini}$. The initialization of coefficient $\hat{\mathbf{z}}_{ini}$ is the multilevel Haar transformation of $\hat{\mathbf{f}}_{ini}$: $\hat{\mathbf{z}}_{ini} = \mathbf{D}^{-1}\hat{\mathbf{f}}_{ini}$. In this article, we choose the level of transformation such that $\mathbf{z}$ has a sparse structure. As shown in Figure 4, when the transform level is small, for example two levels, the coefficient $\mathbf{z}$ is not sparse; when the transform level is sufficiently large, the coefficient is sparse. In this paper, we set $\mathbf{z}$ as a five-level Haar transform coefficient.

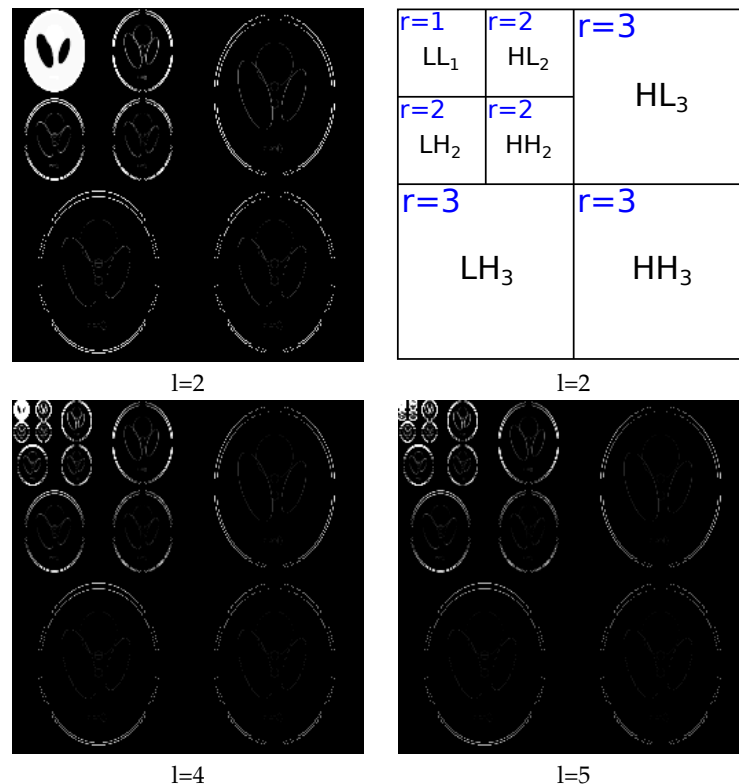


Figure 4. Slice of the 3D multilevel Haar transformation coefficient $\mathbf{z}$ of the 256^3 Shepp-Logan phantom with: 2 levels (top-left), 4 levels (bottom-left), and 5 levels (bottom-right). The figure on the top-right shows the ranks of coefficients for a two-level transformation.

The initialization for α_{z_0} and β_{z_0} is based on the sparse structure of the variable $\mathbf{z}$. In Figure 4, we can see that the sparsity rate depends on the rank of transform coefficient r , where $r \in [1, l + 1]$.

For example, when $l = 2$, shown in Figure 4, the coefficient has three ranks, $r \in [1, 2, 3]$. The first rank $r = 1$ corresponds to the low frequency components in the transform coefficient.

The variable $\mathbf{z}$ is modeled by a St_g distribution, with variance equal to $\text{Var}[z_j | \alpha_{z_0}, \beta_{z_0}] = \beta_{z_0} / (\alpha_{z_0} - 1)$, $\forall j \in [0, N]$. In this article, we fix the value $\alpha_{z_0} = 2.1$ in order to have $\text{Var}[z_j | \alpha_{z_0} = 2.1, \beta_{z_0}] \approx \beta_{z_0}$, $\forall j \in [0, N]$. The sparsity rate of $\mathbf{z}$ is defined by initializing different values for β_{z_0} , with a sparser structure when β_{z_0} is smaller. β_{z_0} is initialized as $\beta_{z_0} = 10^{-r+1}$. When $l = 5$, we have $r = [1, 2, 3, 4, 5, 6]$, and the hyperparameters $\beta_{z_0} = [10^0, 10^{-1}, 10^{-2}, 10^{-3}, 10^{-4}, 10^{-5}]$, respectively.

The initialization for the hyperparameters, α_{ϵ_0} , β_{ϵ_0} , α_{ζ_0} , and β_{ζ_0} , is based on the prior models of the variances v_ϵ and v_ζ we have chosen. In the proposed method, we consider the background of the generalized Student- t distribution, in which both ϵ and ζ are modeled by a Gaussian distribution with inverse Gamma distributed variance, i.e., the St_g distribution according to Equation (13).

The noise ϵ depends on the SNR of the dataset. In order to exploit this information in the initialization, we express the biased dataset as the sum of uncontaminated dataset $\mathbf{g}_0$ and the additive noise ϵ :

$$\mathbf{g} = \mathbf{g}_0 + \epsilon. \quad (36)$$

As the noise ϵ and the uncontaminated data $\mathbf{g}_0$ are supposed to be independent, we have:

$$\|\mathbf{g}\|^2 = \|\mathbf{g}_0\|^2 + \|\epsilon\|^2. \quad (37)$$

The SNR of the dataset is:

$$\text{SNR} = 10 \log \frac{\|\mathbf{g}_0\|^2}{\|\epsilon\|^2} = 10 \log \frac{\|\mathbf{g}\|^2 - \|\epsilon\|^2}{\|\epsilon\|^2}. \quad (38)$$

With $\mathbb{E}[\epsilon] = 0$, we have:

$$v_\epsilon = \mathbb{E}[\epsilon^2] \approx \frac{\|\epsilon\|^2}{M} = \frac{\|\mathbf{g}\|^2}{M} \times \frac{1}{1 + 10^{\text{SNR}/10}}. \quad (39)$$

The mean of variance v_ϵ of the noise ϵ is $\mathbb{E}[v_{\epsilon_i} | \alpha_{\epsilon_0}, \beta_{\epsilon_0}] = \beta_{\epsilon_0} / (\alpha_{\epsilon_0} - 1)$, $\forall i \in [0, M]$, so we obtain:

$$\beta_{\epsilon_0} = \frac{\|\mathbf{g}\|^2}{M} \times \frac{1}{1 + 10^{\text{SNR}/10}} \times (\alpha_{\epsilon_0} - 1). \quad (40)$$

The two hyperparameters α_{ϵ_0} and β_{ϵ_0} are combined according to Equation (40); hence, initialization for one of them is sufficient. In real applications, the SNR of the dataset is unknown, but we can use the projection of an empty object, i.e., $\mathbf{f} = \mathbf{0}$, to obtain a rough value of the variance of noise v_ϵ .

Figure 5 shows the influence of the value of α_{ϵ_0} on the reconstruction. According to the results, a bigger value for α_{ϵ_0} results in a smaller value on RMSE for different numbers of projections and the SNR of the dataset. This monotonous property facilitates the initialization of this hyperparameter, as a large value for α_{ϵ_0} satisfies all cases. When α_{ϵ_0} is greater than a threshold value, the RMSE does not change with different initialization values for α_{ϵ_0} .

For ζ , both α_{ζ_0} and β_{ζ_0} are analyzed for the influence of the reconstruction results.

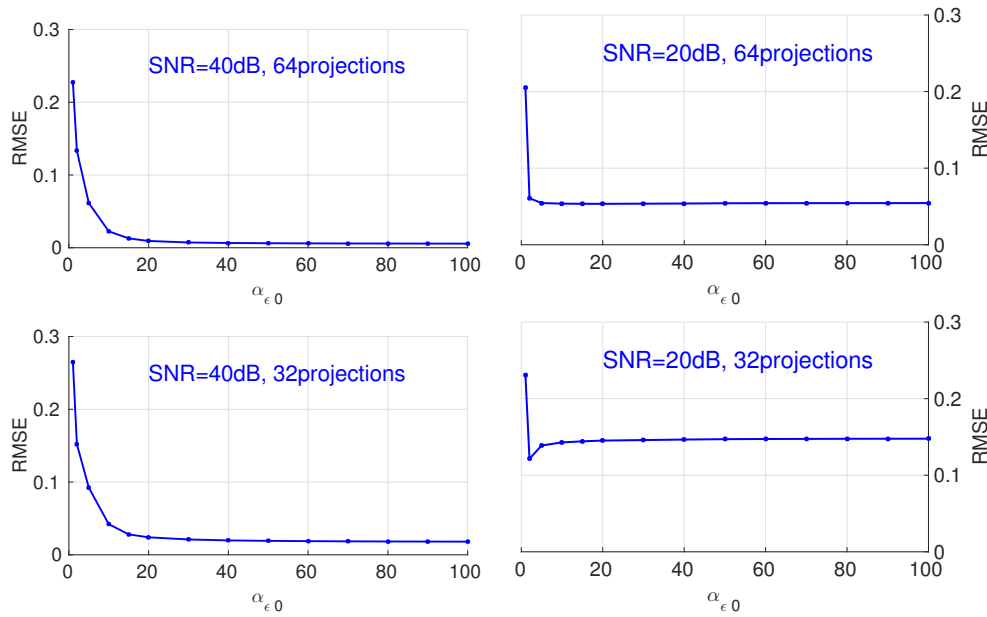


Figure 5. Influence of hyperparameter $\alpha_{\epsilon 0}$ on the RMSE of final reconstruction results for different numbers of projections and noise.

Figure 6 shows the influence of the hyperparameter $\alpha_{\xi 0}$. Different colors represent an initialization with a different value of $\beta_{\xi 0}$. As here we focus on the analysis of $\alpha_{\xi 0}$ for all cases of $\beta_{\xi 0}$ values, we do not show the corresponding $\beta_{\xi 0}$ value for each different color. For different noise levels, different numbers of projections, and different $\beta_{\xi 0}$ values, the RMSE has an upward trend when the value of $\alpha_{\xi 0}$ becomes larger.

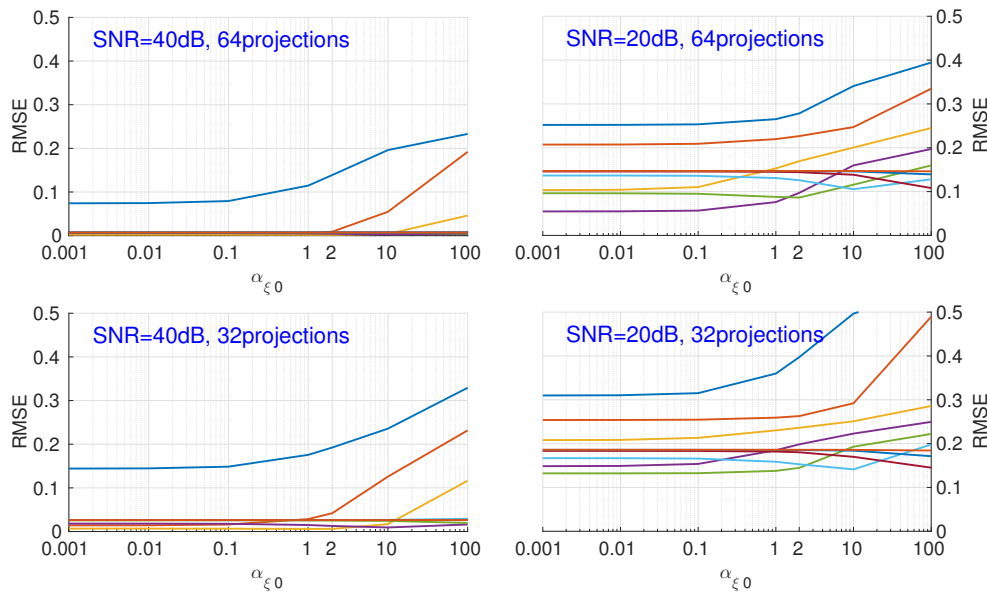


Figure 6. Influence of hyperparameter $\alpha_{\xi 0}$, with different fixed values of $\beta_{\xi 0}$, on the RMSE of reconstruction results for different numbers of projections and noise. Each color corresponds to a different initialization of $\beta_{\xi 0}$.

Figure 7 shows the influence of the hyperparameter $\beta_{\xi 0}$. Different colors represent an initialization with a different hyperparameter $\alpha_{\xi 0}$. For different noise levels, different numbers of projections, and different values for $\alpha_{\xi 0}$, when the value of $\beta_{\xi 0}$ increases, the RMSE decreases, then, after a slight increase, levels out.

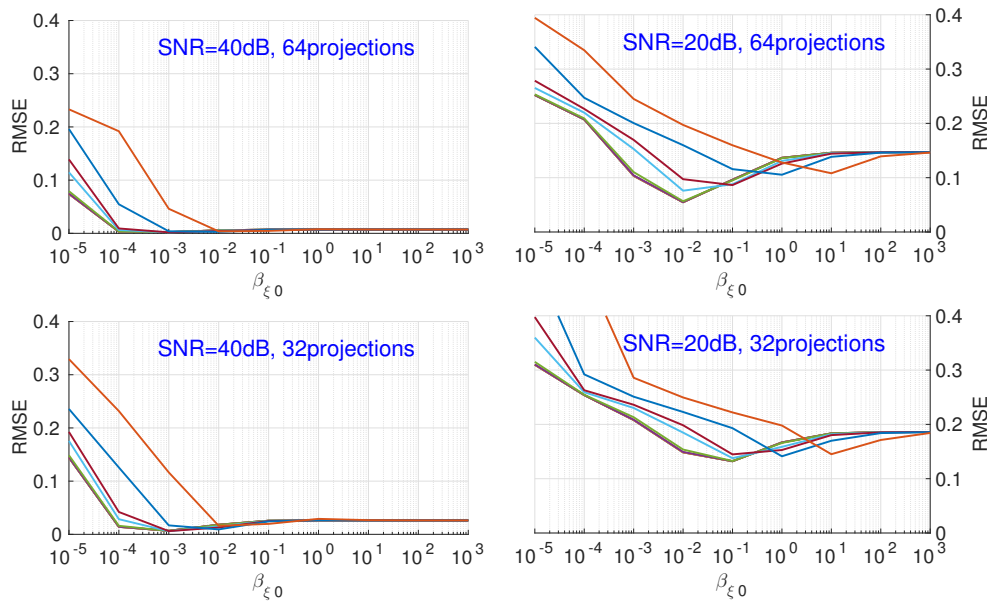


Figure 7. Influence of hyperparameter β_{ξ_0} , with different fixed values of α_{ξ_0} , on the RMSE of reconstruction results for different numbers of projections and noise. Each color corresponds to a different initialization of α_{ξ_0} .

In [46], it is pointed out that when α and β of the St_g distribution are both large, the St_g distribution approaches a Gaussian distribution, which is the case for the additive noise ϵ . If α and β are very small (approaching zero), the St_g distribution becomes a non-informative distribution (Jeffreys distribution); when α and β are both small, the St_g has the sparsity enforcing property, which is the case for the sparse ξ . Consequently, the initialization of the hyperparameters is theoretically supported, and they can be initialized with respect to these properties in other simulations.

3.2. Simulation Results with a Limited Number of Projections

We apply 180, 90, 60, 45, 36, and 18 projections evenly distributed in $[0, 180]$ degrees for the reconstruction of the 3D Shepp–Logan phantom of size 256^3 ; each projection contains 256^2 detectors. The number of projections is chosen such that there is respectively one projection every 1, 2, 3, 4, 5, and 10 degrees.

In Table 1, different evaluation metrics of the reconstructed 256^3 Shepp–Logan phantom are compared. It is shown that the HHBM method does not always perform better than the TV method, especially when there are sufficient numbers of projections. However, when there is insufficient projection data, the HHBM method is more robust than the TV method. On the other hand, as it is known that the choice of regularization parameter plays an important role in the regularization methods like QR or TV and the value for the regularization parameter should be selected for each different system settings, the HHBM method is much more robust on the initialization of hyper-parameters. As we can see from Figures 5–7, once we have chosen the hyperparameters in a certain interval, which is not difficult to fix according to the properties of the prior model, we can obtain the appropriate reconstruction results. More importantly, in the Bayesian approach, the prior model can be chosen from a variety of other suitable distributions, which gives more possibilities for the models than the conventional regularization methods. We may also choose different point estimators from the posterior distribution, for example the posterior mean, etc.

Table 1. RMSE, ISNR, PSNR, and SSIM of the reconstructed phantom with 50 global iterations (10 gradient descent iterations in each global iteration). The values of the regularization parameters are respectively $\lambda_{QR} = 10$ and $\lambda_{TV} = 50$ for SNR = 40 dB, $\lambda_{QR} = 600$ and $\lambda_{TV} = 100$ for SNR = 20 dB. TV, Total Variation; HHBM, Hierarchical Haar transform-based Bayesian Method; QR, Quadratic Regularization.

256 × 256 × 256 Shepp–Logan Phantom												
	180 Projections						90 Projections					
	40 dB			20 dB			40 dB			20 dB		
	QR	TV	HHBM	QR	TV	HHBM	QR	TV	HHBM	QR	TV	HHBM
RMSE	0.0236	0.0114	0.0069	0.1309	0.0209	0.0755	0.0401	0.0212	0.0092	0.1558	0.0491	0.1117
ISNR	5.5584	8.7217	10.9346	7.2024	15.1775	10.2162	6.6136	9.3832	12.9973	8.4583	13.4765	9.9056
PSNR	30.0675	33.2308	35.4437	22.6318	30.6069	25.0209	27.7743	30.5439	34.1579	21.8754	26.8937	23.3227
SSIM	0.9999	0.9999	1.0000	0.9992	0.9999	0.9995	0.9997	0.9999	0.9999	0.9990	0.9997	0.9993
	60 Projections						45 Projections					
	40 dB			20 dB			40 dB			20 dB		
	QR	TV	HHBM	QR	TV	HHBM	QR	TV	HHBM	QR	TV	HHBM
RMSE	0.0636	0.0321	0.0107	0.1656	0.0753	0.1293	0.0904	0.0474	0.0132	0.1854	0.0901	0.1414
ISNR	9.3826	12.3480	17.1346	9.1492	12.5701	10.2226	10.3301	13.1308	18.6839	10.0137	13.1476	11.1916
PSNR	25.7693	28.7347	33.5214	21.6116	25.0325	22.6849	24.2404	27.0412	32.5942	21.1195	24.2535	22.2974
SSIM	0.9996	0.9995	0.9999	0.9990	0.9995	0.9992	0.9994	0.9997	0.9999	0.9988	0.9994	0.9991
	36 Projections						18 Projections					
	40 dB			20 dB			40 dB			20 dB		
	QR	TV	HHBM	QR	TV	HHBM	QR	TV	HHBM	QR	TV	HHBM
RMSE	0.1177	0.0680	0.0169	0.1957	0.1116	0.1500	0.2581	0.2104	0.0574	0.2907	0.2313	0.2014
ISNR	10.6591	13.0424	19.0933	10.8633	13.3032	12.0187	10.7122	11.5992	17.2373	10.8088	11.8022	12.4036
PSNR	23.0949	25.4783	31.5292	20.8865	23.3264	22.0420	19.6263	20.5133	26.1514	19.1085	20.1020	20.7033
SSIM	0.9993	0.9996	0.9999	0.9988	0.9993	0.9990	0.9983	0.9987	0.9996	0.9981	0.9985	0.9987

Figures 8 and 9 show the reconstructed middle slice of the “Shepp–Logan” phantom and “head” object by using the TV and HHBM methods from 36 projections with SNR = 40 dB and SNR = 20 dB. The red curve illustrates the profile of the blue line position. In the reconstructed Shepp–Logan phantom obtained using the TV method, the three small circles on the top of the slice are not evident. By using the HHBM method, we can distinguish these three small circles. By comparing the profiles of the slice of the reconstructed Shepp–Logan phantom, we can see that by using the HHBM method, the contour positions on the profile are closer to the original profile than those obtained using the TV method. In the reconstructed head object, there are more details than the simulated Shepp–Logan phantom, especially in the zoom area in the second line in Figure 8. By comparing the results, we can see that for the type of object that contains some small details, the TV method derives a result with smoother homogeneous areas, but with fewer details in the contour areas than the HHBM method. Some of the white material, which is dispersed into discontinuous small blocks in the head object, is connected in the results of the TV method. From these images, we conclude that with an insufficient number of projections, the proposed method gives results with clearer contours and details.

Figure 10 shows the reconstructed Shepp–Logan phantom from 18 projections of SNR = 40 dB and 20 dB. In this very underdetermined case, the HHBM method can still obtain a result that is clear enough to distinguish the primary zones and contours of the object.

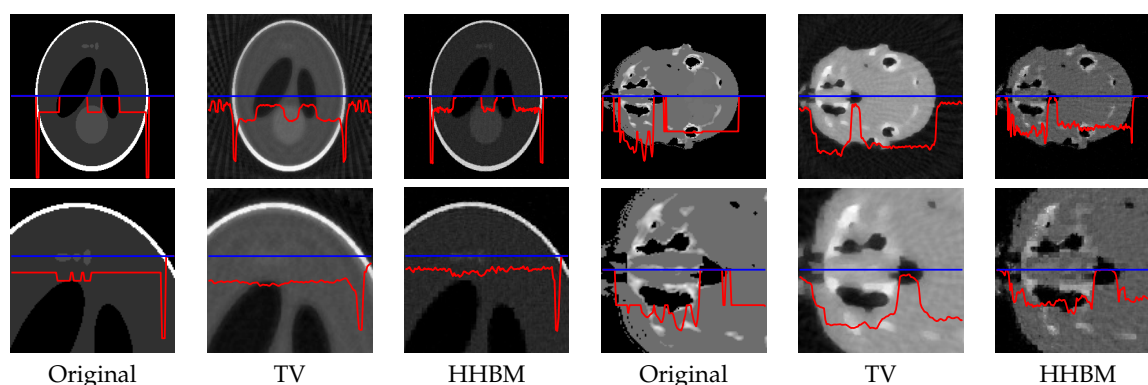


Figure 8. Reconstructed “Shepp–Logan” phantom and “head” object of size 256^3 , with a dataset of 36 projections and SNR = 40 dB, by using the TV and HHBM methods. The bottom figures are zones of the corresponding top figures. The red curves are the profiles at the position of the corresponding blue lines.

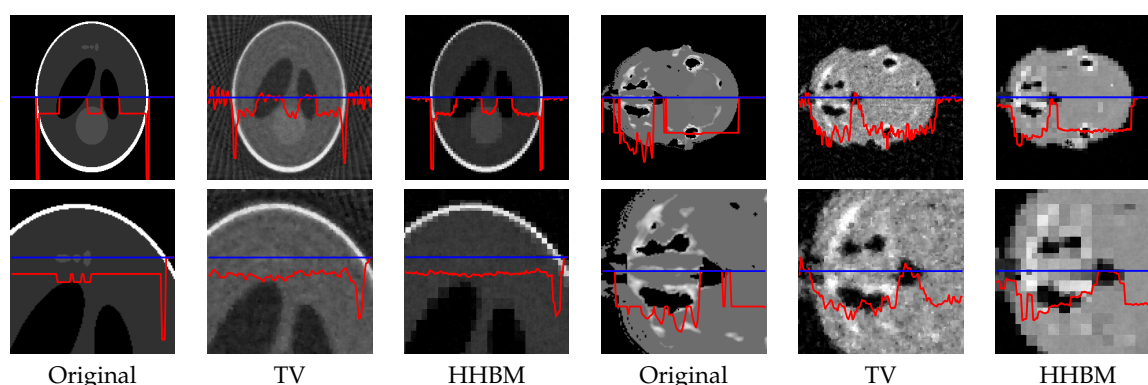


Figure 9. Reconstructed “Shepp–Logan” phantom and “head” object of size 256^3 , with a dataset of 36 projections and SNR = 20 dB, by using the TV and HHBM methods. Bottom figures are zones of the corresponding top figures. The red curves are the profiles at the position of the corresponding blue lines.

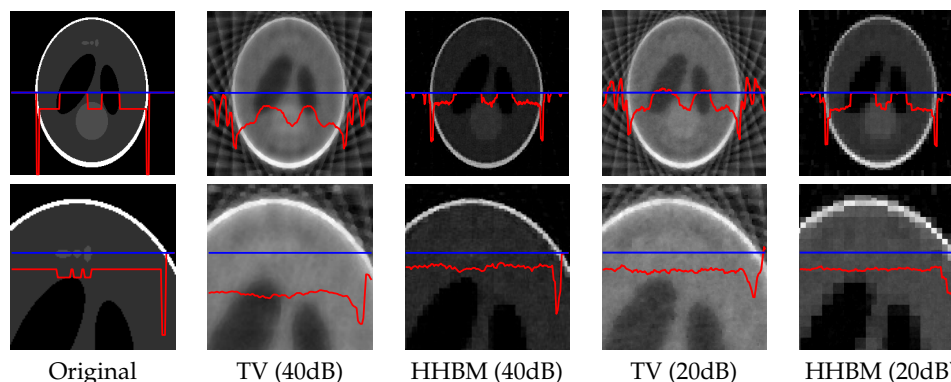


Figure 10. Reconstructed Shepp–Logan phantom of size 256^3 , with a dataset of 18 projections of SNR = 40 dB and 20 dB, by using the TV and HHBM methods. The red curves are the profiles at the position of the corresponding blue lines.

Figures 11 and 12 show the comparison between the QR, TV, and HHBM methods with a high SNR = 40 dB and a low SNR = 20 dB dataset, respectively. The abscissa corresponds to the number of projections evenly distributed from 0° – 180° , and the ordinate is the RMSE after 50 iterations. In the simulations, we used an SNR = 40 dB to represent a weak noise case and SNR = 20 dB for a strong noise case. When SNR = 40 dB, the HHBM method outperforms both quadratic regularization and the TV method. When SNR = 20 dB, the TV method with the optimal regularization parameter outperforms the HHBM method. However, in Figure 12, we show another curve in light green (TV2) showing the TV reconstruction with some random regularization parameters, which are chosen as a value not far from the optimal regularization parameters. From these results, we can see that the HHBM method is more robust than the TV reconstruction method with respect to the regularization parameter, the optimal value of which is, on the other hand, difficult to determine in the real applications where we cannot evaluate the estimation quality.

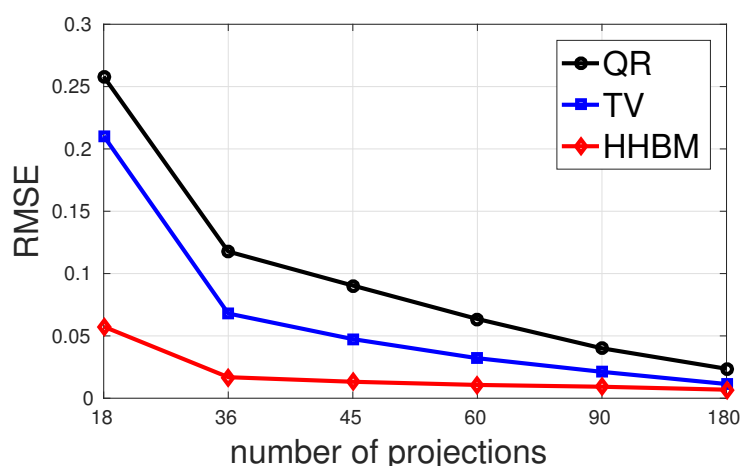


Figure 11. The performances of different methods for reconstructing the Shepp–Logan phantom in terms of RMSE with different numbers of projections evenly distributed in $[0^\circ, 180^\circ]$ and a high SNR = 40 dB.

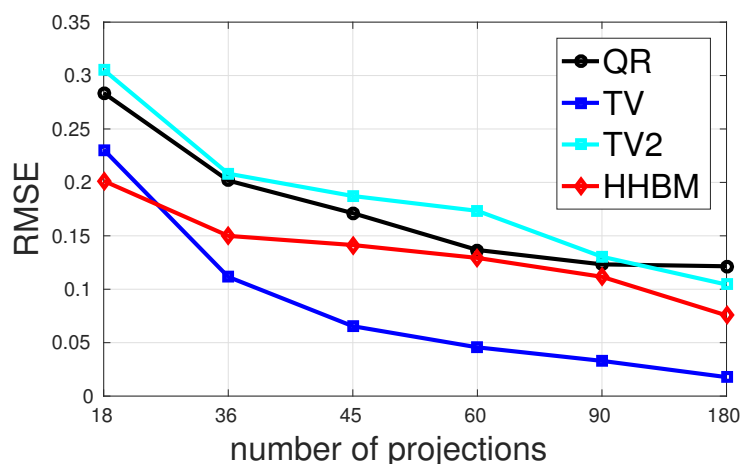


Figure 12. The performances of different methods for reconstructing the Shepp–Logan phantom in terms of RMSE with different numbers of projections evenly distributed in $[0^\circ, 180^\circ]$ and a low SNR = 20 dB.

3.3. Simulation Results with a Limited Angle of Projections

In both medical and industrial X-ray CT, another common challenge is the limit of projection angles. In this part of the simulation, we use evenly-distributed projections in a limited range of angles for the simulated 3D “Shepp–Logan” phantom and the 3D “head” object, both of which have a size of 256^3 .

Figures 13 and 14 show the middle slice of the reconstructed Shepp–Logan phantom and the head object, from 90 projections distributed between 0° and 90° , with projection SNR of 40 dB and 20 dB. By using the TV method, the reconstructed object is blurry along the diagonal direction for which there is no projection data, and there is a square corner where the object should have a rounded edge. By using the HHBM method, we get results that are more consistent with the true shape and clearer contours.

Figures 15 and 16 show the comparison of the performance in terms of RMSE of different methods with a high SNR = 40 dB and a low SNR = 20 dB. In this comparison, four cases of limited projection angles are considered, and they are: 45, 90, 135, and 180 projections evenly distributed in $[0^\circ, 45^\circ]$, $[0^\circ, 90^\circ]$, $[0^\circ, 135^\circ]$, and $[0^\circ, 180^\circ]$, respectively. There is one projection every 1° . From these two figures, we conclude that the proposed HHBM method remains more robust than the other two conventional methods when there are limited projection angles.

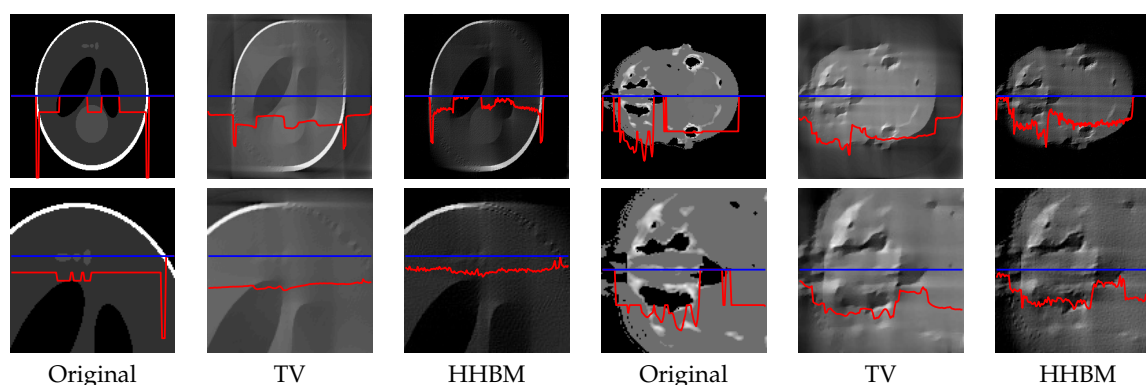


Figure 13. Slice of the reconstructed 3D Shepp–Logan phantom and 3D head object, with 90 projections evenly distributed in $[0^\circ, 90^\circ]$, SNR = 40 dB. The bottom figures are parts of the corresponding top figures. The red curves are the profiles at the position of the corresponding blue lines.

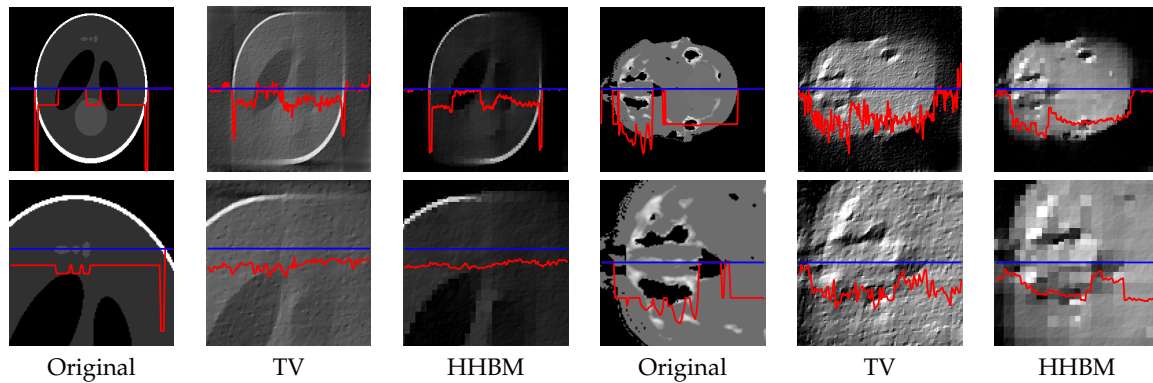


Figure 14. Slice of the reconstructed 3D Shepp–Logan phantom and 3D head object, with 90 projections evenly distributed in $[0^\circ, 90^\circ]$, SNR = 20 dB. Bottom figures are parts of the corresponding top figures. The red curves are the profiles at the position of the corresponding blue lines.

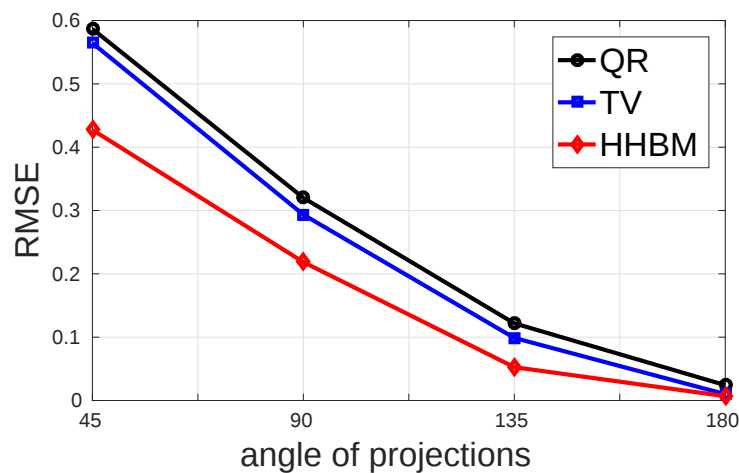


Figure 15. The performance of different methods for reconstructing the Shepp–Logan phantom in terms of RMSE with different limited projection angles and a high SNR = 40 dB.

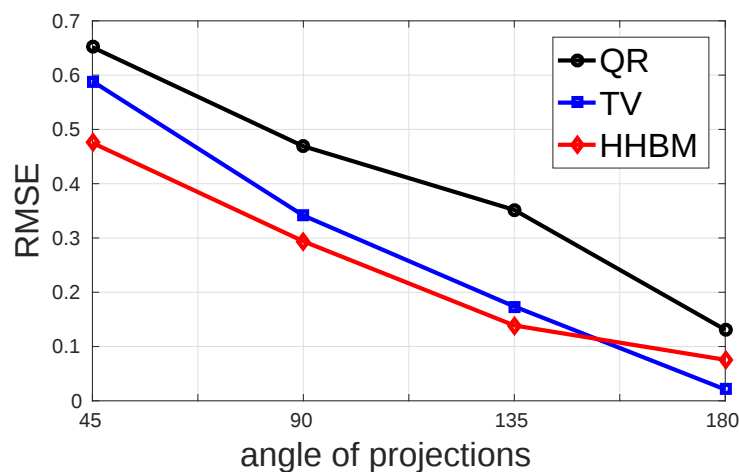


Figure 16. The performance of different methods for reconstructing the Shepp–Logan phantom in terms of RMSE with different limited projection angles and a low SNR = 20 dB.

3.4. Simulation with a Different Forward Model

To study the robustness of the results with respect to the modeling errors, we apply a slightly different forward model for the one used for the generation of the simulated data and the one used for the reconstruction. During the projection, the projector is applied on the Shepp–Logan phantom of size 1024^3 , with the detector size of 256^2 for each projection direction.

In Figure 17, we show the results of the two projection data obtained, as well as their difference δ_{proj} , which can be considered as the forward modeling error. The δ_{proj} is rescaled in order to show clearly the details. We then use the data obtained from the 1024^3 phantom to reconstruct an object of size 256^3 .

In Figure 18, the middle slices of the reconstructed Shepp–Logan phantom by using the QR, TV, and HHBM methods are presented, by using 180 and 36 projections, respectively. From the figures, we can see that when there are 180 projections, all three methods performs well, and the TV method detects better the contours, while the HHBM method has more noise at the contour areas. When there are insufficient projection numbers (36 projections in this simulation), the HHBM method outperforms the QR and TV methods for reconstructing the details in the phantom, for example the three small circles in the top of the phantom.

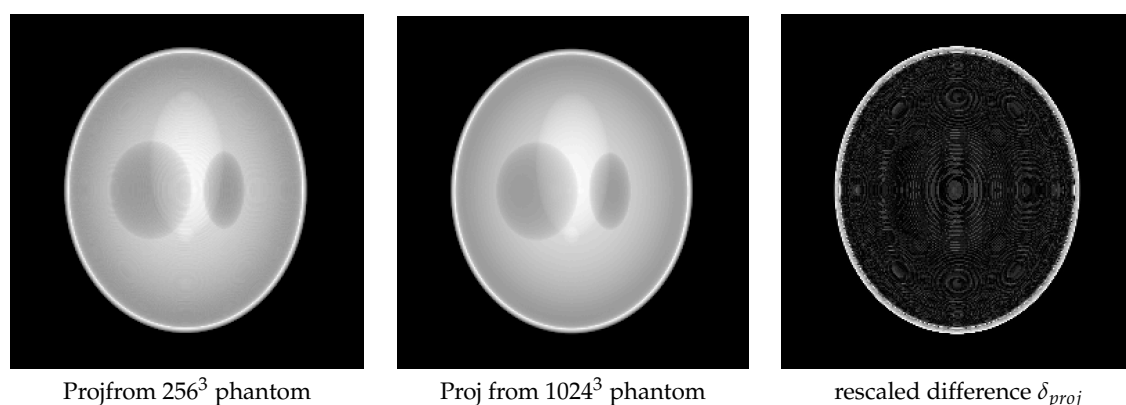


Figure 17. The projection image at an angle of 90 degrees by using two different forward models: projection from the 256^3 phantom (left) and projection from the 1024^3 phantom (middle). The difference between them is shown on the right.

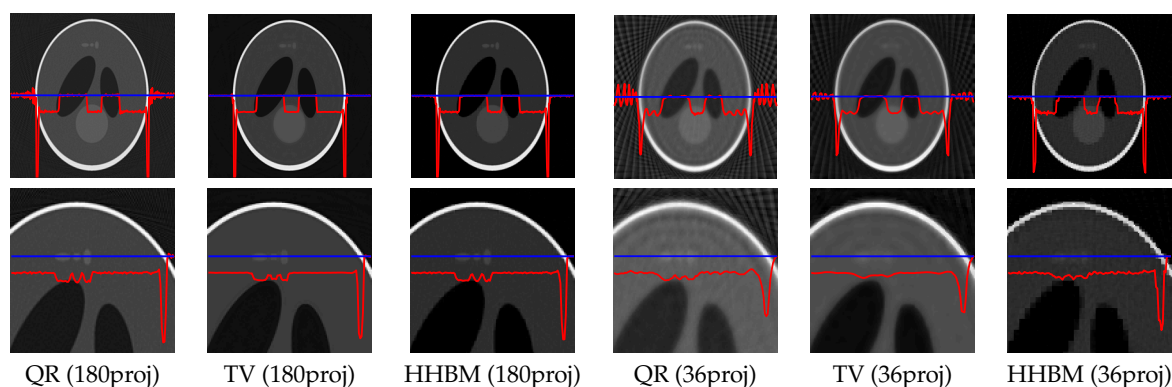


Figure 18. Reconstructed phantom with different forward models and 180 projections and 36 projections by using the QR (left), TV (middle), and HHBM (right) methods.

In Table 2, the RMSE of the reconstructed phantom by using the different methods are compared. We can conclude that, when the projector model is different than the reconstruction one, all these three methods (QR, TV, and HHBM) have good performance when there are 180 projections. When the projection number decreases, the TV and HHBM methods outperform the QR method. Comparing with the TV method, the HHBM method is more robust to the number of projections. When the projection number is smaller than 60, the HHBM method outperforms the TV method.

All the MATLAB codes for the simulations in this paper can be found on GitHub [64].

Table 2. The RMSE of the reconstructed 256^3 Shepp–Logan phantom by using projection obtained from a 1024^3 Shepp–Logan phantom.

RMSE	QR	TV	HHBM
180 proj	0.0581	0.0540	0.0558
90 proj	0.0655	0.0573	0.0610
60 proj	0.0846	0.0675	0.0690
45 proj	0.1079	0.0830	0.0783
36 proj	0.1326	0.1027	0.0882

4. Discussion

One advantage of the Bayesian approach is the estimation of the parameters along with the estimation of unknown variables of the forward model at each iteration. However, like in regularization methods, the hyper-parameters need to be initialized.

While the parameters in the regularization methods play an important role in the final results and they are costly to fix, the hyper-parameters in HHBM can be initialized based on the prior information (the sparse structure of z) and the prior model (the Student- t distribution). In this article, we have shown that once the hyper-parameters are fixed in a certain appropriate interval, which is not difficult to obtain, the corresponding algorithm is robust. In this work, the hyper-parameters are not fixed via the classical approach using non-informative prior laws (i.e., considering the inverse Gamma corresponding parameters such that they approach Jeffreys') [65].

5. Conclusions

In this paper, we propose a Bayesian method with a hierarchical structured prior model based on multilevel Haar transformation (HHBM) for 3D X-ray CT reconstruction. Simulation results indicate that for a limited number of projections or limited projection angles, the proposed method is more robust to noise and to regularization parameters than the classical QR and TV methods.

Indeed, we observe a relatively weak influence of the hyper-parameters in the behavior of the corresponding iterative algorithm. The interest of this weak dependency is that it offers a practical way to ensure the initialization of the algorithm, which typically is not-trivial.

In this context, as future work, we are investigating the causes of the relatively weak influence of the hyper-parameters and the theoretical foundation of the corresponding robust interval, extending the discussion to the same approach using sparsity-enforcing priors expressed as normal variance mixtures, but for other mixing distributions (Gamma, generalized inverse Gaussian) [66].

Another extension of this work is to consider the posterior mean as an estimator. This can be done via the Variational Bayesian Approach (VBA), but a practical implementation requires a method of accessing the diagonal elements of the large matrix $H^T H$, which is being studied by our group.

Author Contributions: This work has been done during the PhD preparation of the first author. A.M.-D. and N.G. have been the co-supervisors of the PhD work. M.D. was a post-doc who also had finished his PhD under the supervision of A.M.-D. and was continuing to work on this project.

Funding: This research was funded by China Scholarship Council.

Conflicts of Interest: The authors declare no conflict of interest.

References

1. Arridge, S.R. Optical tomography in medical imaging. *Inverse Probl.* **1999**, *15*, R41, doi:10.1088/0266-5611/15/2/022.
2. Hanke, R.; Fuchs, T.; Uhlmann, N. X-ray based methods for non-destructive testing and material characterization. *Nucl. Instrum. Methods Phys. Res. Sect. A* **2008**, *591*, 14–18.
3. Kalender, W.A. X-ray Computed Tomography. *Phys. Med. Biol.* **2006**, *51*, R29.

4. Natterer, F.; Wubbeling, F. A propagation-backpropagation method for ultrasound tomography. *Inverse Probl.* **1995**, *11*, 1225, doi:10.1088/0266-5611/11/6/007.
5. Bolomey, J.C.; Pichot, C. Microwave tomography: from theory to practical imaging systems. *Int. J. Imaging Syst. Technol.* **1990**, *2*, 144–156.
6. Nelson, J.; Milner, T.; Tanenbaum, B.; Goodman, D.; Van Gemert, M. Infra-red tomography of port-wine-stain blood vessels in human skin. *Lasers Med. Sci.* **1996**, *11*, 199–204.
7. Deans, S.R. *The Radon Transform and Some of Its Applications*; Courier Corporation: Mineola, NY, USA, 2007.
8. Feldkamp, L.; Davis, L.; Kress, J. Practical cone-beam algorithm. *JOSA A* **1984**, *1*, 612–619.
9. Alifanov, O.M.; Artioukhine, E.A.; Rummyantsev, S.V. *Extreme Methods for Solving Ill-Posed Problems with Applications to Inverse Heat Transfer Problems*; Begell House: New York, NY, USA, 1995.
10. O'Sullivan, F. A statistical perspective on ill-posed inverse problems. *Stat. Sci.* **1986**, *1*, 502–518.
11. Gordon, R. A tutorial on ART (algebraic reconstruction techniques). *IEEE Trans. Nucl. Sci.* **1974**, *21*, 78–93.
12. Gordon, R.; Bender, R.; Herman, G.T. Algebraic reconstruction techniques (ART) for three-dimensional electron microscopy and X-ray photography. *J. Theor. Biol.* **1970**, *29*, 471–481.
13. Andersen, A.H.; Kak, A.C. Simultaneous algebraic reconstruction technique (SART): A superior implementation of the ART algorithm. *Ultrason. Imaging* **1984**, *6*, 81–94.
14. Trampert, J.; Leveque, J.J. Simultaneous iterative reconstruction technique: physical interpretation based on the generalized least squares solution. *J. Geophys. Res.* **1990**, *95*, 553–559.
15. Bangliang, S.; Yiheng, Z.; Lihui, P.; Danya, Y.; Baofen, Z. The use of simultaneous iterative reconstruction technique for electrical capacitance tomography. *Chem. Eng. J.* **2000**, *77*, 37–41.
16. Coric, S.; Leeser, M.; Miller, E.; Trepanier, M. Parallel-beam backprojection: An FPGA implementation optimized for medical imaging. In Proceedings of the 2002 ACM/SIGDA Tenth International Symposium on Field-Programmable Gate Arrays, Monterey, CA, USA, 24–26 February 2002; pp. 217–226.
17. Ay, M.R.; Zaidi, H. Development and validation of MCNP4C-based Monte Carlo simulator for fan-and cone-beam x-ray CT. *Phys. Med. Biol.* **2005**, *50*, 4863–4885.
18. Mozzo, P.; Procacci, C.; Tacconi, A.; Tinazzi Martini, P.; Bergamo Andreis, I. A new volumetric CT machine for dental imaging based on the cone-beam technique: preliminary results. *Eur. Radiol.* **1998**, *8*, 1558–1564.
19. Ben-Israel, A.; Greville, T.N. *Generalized Inverses: Theory and Applications*; Springer: New York, NY, USA, 2003; Volume 15.
20. Calvetti, D.; Morigi, S.; Reichel, L.; Sgallari, F. Tikhonov regularization and the L-curve for large discrete ill-posed problems. *J. Comput. Appl. Math.* **2000**, *123*, 423–446.
21. Chambolle, A.; Lions, P.L. Image recovery via total variation minimization and related problems. *Numer. Math.* **1997**, *76*, 167–188.
22. Goldstein, T.; Osher, S. The split Bregman method for L1-regularized problems. *SIAM J. Imaging Sci.* **2009**, *2*, 323–343.
23. Sidky, E.Y.; Pan, X. Image reconstruction in circular cone-beam computed tomography by constrained, total-variation minimization. *Phys. Med. Biol.* **2008**, *53*, 4777, doi:10.1088/0031-9155/53/17/021.
24. Chan, T.F.; Golub, G.H.; Mulet, P. A nonlinear primal-dual method for total variation-based image restoration. *SIAM J. Sci. Comput.* **1999**, *20*, 1964–1977.
25. Wahlberg, B.; Boyd, S.; Annergren, M.; Wang, Y. An ADMM algorithm for a class of total variation regularized estimation problems. *IFAC Proc. Vol.* **2012**, *45*, 83–88.
26. Zou, H.; Hastie, T. Regularization and variable selection via the elastic net. *J. R. Stat. Soc. Ser. B* **2005**, *67*, 301–320.
27. Ramani, S.; Liu, Z.; Rosen, J.; Nielsen, J.F.; Fessler, J.A. Regularization parameter selection for nonlinear iterative image restoration and MRI reconstruction using GCV and SURE-based methods. *IEEE Trans. Image Process.* **2012**, *21*, 3659–3672.
28. Galatsanos, N.P.; Katsaggelos, A.K. Methods for choosing the regularization parameter and estimating the noise variance in image restoration and their relation. *IEEE Trans. Image Process.* **1992**, *1*, 322–336.
29. Wang, G.; Schultz, L.; Qi, J. Statistical image reconstruction for muon tomography using a Gaussian scale mixture model. *IEEE Trans. Nucl. Sci.* **2009**, *56*, 2480–2486.
30. Redner, R.A.; Walker, H.F. Mixture densities, maximum likelihood and the EM algorithm. *SIAM Rev.* **1984**, *26*, 195–239.
31. Moon, T.K. The expectation-maximization algorithm. *IEEE Signal Process. Mag.* **1996**, *13*, 47–60.

32. Tregouet, D.; Escolano, S.; Turet, L.; Mallet, A.; Golmard, J. A new algorithm for haplotype-based association analysis: the Stochastic-EM algorithm. *Ann. Hum. Genet.* **2004**, *68*, 165–177.
33. Hudson, H.M.; Larkin, R.S. Accelerated image reconstruction using ordered subsets of projection data. *IEEE Trans. Med. Imaging* **1994**, *13*, 601–609.
34. Gelman, A.; Carlin, J.B.; Stern, H.S.; Rubin, D.B. *Bayesian Data Analysis*; Chapman & Hall/CRC: Boca Raton, FL, USA, 2014; Volume 2.
35. Bali, N.; Mohammad-Djafari, A. Bayesian approach with hidden Markov modeling and mean field approximation for hyperspectral data analysis. *IEEE Trans. Image Process.* **2008**, *17*, 217–225.
36. Mohammad-Djafari, A. Joint estimation of parameters and hyperparameters in a Bayesian approach of solving inverse problems. In Proceedings of the 3rd IEEE International Conference on Image Processing, Lausanne, Switzerland, 19 September 1996; Volume 2, pp. 473–476.
37. Kolehmainen, V.; Vanne, A.; Siltanen, S.; Jarvenpaa, S.; Kaipio, J.P.; Lassas, M.; Kalke, M. Parallelized Bayesian inversion for three-dimensional dental X-ray imaging. *IEEE Trans. Med. Imaging* **2006**, *25*, 218–228.
38. Qi, J.; Leahy, R.M. Resolution and noise properties of MAP reconstruction for fully 3-D PET. *IEEE Trans. Med. Imaging* **2000**, *19*, 493–506.
39. Ericson, W.A. A note on the posterior mean of a population mean. *J. R. Stat. Soc. Ser. B* **1969**, *31*, 332–334.
40. Fox, C.W.; Roberts, S.J. A tutorial on variational Bayesian inference. *Artif. Intell. Rev.* **2012**, *38*, 85–95.
41. Tzikas, D.G.; Likas, A.C.; Galatsanos, N.P. The variational approximation for Bayesian inference. *IEEE Signal Process. Mag.* **2008**, *25*, 131–146.
42. Ayasso, H.; Mohammad-Djafari, A. Joint NDT image restoration and segmentation using Gauss–Markov–Potts prior models and variational Bayesian computation. *IEEE Trans. Image Process.* **2010**, *19*, 2265–2277.
43. Noël, P.B.; Walczak, A.M.; Xu, J.; Corso, J.J.; Hoffmann, K.R.; Schafer, S. GPU-based cone beam computed tomography. *Comput. Methods Programs Biomed.* **2010**, *98*, 271–277.
44. Gac, N.; Mancini, S.; Desvignes, M.; Houzet, D. High speed 3D tomography on CPU, GPU, and FPGA. *EURASIP J. Embed. Syst.* **2008**, *2008*, 930250, doi:10.1155/2008/930250.
45. Wang, L.; Mohammad-Djafari, A.; Gac, N.; Dumitru, M. Computed tomography reconstruction based on a hierarchical model and variational Bayesian method. In Proceedings of the 2016 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP), Shanghai, China, 20–25 March 2016; pp. 883–887.
46. Dumitru, M. A Bayesian Approach for Periodic Components Estimation for Chronobiological Signals. Ph.D. Thesis, Paris Saclay: Paris, France, 2016.
47. Wang, L.; Mohammad-Djafari, A.; Gac, N. X-ray Computed Tomography using a sparsity enforcing prior model based on Haar transformation in a Bayesian framework. *Fundam. Inform.* **2017**, *155*, 449–480.
48. Stanković, R.S.; Falkowski, B.J. The Haar wavelet transform: Its status and achievements. *Comput. Electr. Eng.* **2003**, *29*, 25–44.
49. Starck, J.L.; Candès, E.J.; Donoho, D.L. The curvelet transform for image denoising. *IEEE Trans. Image Process.* **2002**, *11*, 670–684.
50. Do, M.N.; Vetterli, M. The contourlet transform: an efficient directional multiresolution image representation. *IEEE Trans. Image Process.* **2005**, *14*, 2091–2106.
51. Selesnick, I.W.; Baraniuk, R.G.; Kingsbury, N.C. The dual-tree complex wavelet transform. *IEEE Signal Process. Mag.* **2005**, *22*, 123–151.
52. Nadarajah, S. A generalized normal distribution. *J. Appl. Stat.* **2005**, *32*, 685–694.
53. Rasmussen, C.E. The infinite Gaussian mixture model. *NIPS* **1999**, *12*, 554–560.
54. Klebanov, L.B. *Heavy Tailed Distributions*; MatfFyzPress: Prague, Czech Republic, 2003.
55. Frese, T.; Bouman, C.A.; Sauer, K. Adaptive wavelet graph model for Bayesian tomographic reconstruction. *IEEE Trans. Image Process.* **2002**, *11*, 756–770.
56. Delaney, A.H.; Bresler, Y. Multiresolution tomographic reconstruction using wavelets. *IEEE Trans. Image Process.* **1995**, *4*, 799–813.
57. Rantala, M.; Vanska, S.; Jarvenpaa, S.; Kalke, M.; Lassas, M.; Moberg, J.; Siltanen, S. Wavelet-based reconstruction for limited-angle X-ray tomography. *IEEE Trans. Med. Imaging* **2006**, *25*, 210–217.
58. Loris, I.; Nolet, G.; Daubechies, I.; Dahlen, F. Tomographic inversion using l1-norm regularization of wavelet coefficients. *Geophys. J. Int.* **2007**, *170*, 359–370.

59. Sauer, K.; Bouman, C. A local update strategy for iterative reconstruction from projections. *IEEE Trans. Signal Process.* **1993**, *41*, 534–548.
60. Elad, M.; Milanfar, P.; Rubinstein, R. Analysis versus synthesis in signal priors. *Inverse Probl.* **2007**, *23*, 947–968.
61. Boyd, S.; Vandenberghe, L. *Convex Optimization*; Cambridge University Press: Cambridge, UK, 2004.
62. Wang, Z.; Bovik, A.C.; Sheikh, H.R.; Simoncelli, E.P. Image quality assessment: from error visibility to structural similarity. *IEEE Trans. Image Process.* **2004**, *13*, 600–612.
63. Palenstijn, W.J.; Batenburg, K.J.; Sijbers, J. The ASTRA tomography toolbox. In Proceedings of the 13th International Conference on Computational and Mathematical Methods in Science and Engineering, CMMSE, Almeria, Spain, 24–27 June 2013; Volume 2013.
64. Dumitru, M.; Wang, L.; Gac, N.; Mohammad-Djafari, A. iterTomoGPI-GPL, Gitlab Repository, April 2018. Available online: https://github.com/nicolasgac/TomoGPI_for_Astra (accessed on April 2018).
65. Figueiredo, M.A.; Nowak, R.D. Wavelet-based image estimation: An empirical Bayes approach using Jeffrey’s noninformative prior. *IEEE Trans. Image Process.* **2001**, *10*, 1322–1331.
66. Dumitru, M.; Wang, L.; Gac, N.; Mohammad-Djafari, A. Performance Comparison of Bayesian Iterative Algorithms for Three Classes of Sparsity Enforcing Priors With Application in Computed Tomography. In Proceedings of the 2017 IEEE International Conference on Image Processing (ICIP), Beijing, China, 17–20 September 2017.



© 2018 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).