



Bayesian X-ray Computed Tomography using a Three-level Hierarchical Prior Model

Li Wang, Ali Mohammad-Djafari, Nicolas Gac

► To cite this version:

Li Wang, Ali Mohammad-Djafari, Nicolas Gac. Bayesian X-ray Computed Tomography using a Three-level Hierarchical Prior Model. AIP Conference, Bayesian inference and maximum entropy methods in science and engineering (Maxent 2016), Jul 2016, Gent, Belgium. 10.1063/1.4985361 . hal-01403790

HAL Id: hal-01403790

<https://hal.science/hal-01403790>

Submitted on 27 Nov 2016

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.

Bayesian X-ray Computed Tomography using a Three-level Hierarchical Prior Model

Li Wang^{1,2,a)}, Ali Mohammad-Djafari^{1,2,3,b)} and Nicolas Gac^{1,2,c)}

¹*Laboratoire des Signaux et Systèmes, CentraleSupélec, France*

²*Université Paris Sud, France*

³*CNRS, France*

^{a)}Corresponding author: li.wang@lss.supelec.fr

^{b)}djafari@lss.supelec.fr

^{c)}nicolas.gac@lss.supelec.fr

Abstract. In recent decades X-ray Computed Tomography (CT) image reconstruction has been largely developed in both medical and industrial domain. In this paper, we propose using the Bayesian inference approach with a new hierarchical prior model. In the proposed model, a generalised Student-t distribution is used to enforce the Haar transformation of images to be sparse. Comparisons with some state of the art methods are presented. It is shown that by using the proposed model, the sparsity of the sparse representation of images is enforced, so that edges of images are preserved. Simulation results are also provided to demonstrate the effectiveness of the new hierarchical model for reconstruction with fewer projections.

Keywords: Computed Tomography (CT), Bayesian Method, Hierarchical Model, Sparsity, Student-t distribution, Inverse Problem, Variational Bayesian Approach (VBA), Joint Maximum A Posterior (JMAP), Infinite Gaussian Scaled Mixture (IGSM)

INTRODUCTION

Computed Tomography (CT) has been a widely used technique in recent decades. Research on X-ray CT has been developing fast in both medical and industrial applications. By 1990, the famous Filtered Back Projection (FBP) had been developed. In the X-ray CT field, research started to move into the fully 3D domain, as medical and industrial scanners turned to 3D reconstruction methods and iterative reconstruction schemes.

In X ray CT, the intensity of X ray is attenuated when passing through the object, and the parameters to be reconstructed is the linear attenuation coefficient inside the object under the test. The relation between the measured quantity g and the unknown distribution of the attenuation coefficient $f(x, y)$ is given by:

$$g = -\ln \frac{I}{I_0} = \int f(x, y) dl \quad (1)$$

where (x, y) is the coordinate position of the object, I_0 the initial X ray intensity, I the detected X ray intensity after passing through the object and dl the infinite signal length on the line joining the emitter and the detector.

Considering that different tissue has different attenuation coefficient values, for example the metal and air material in a component of industry applications, one can therefore distinguish the structure information through the attenuated data. Often, in order to distinguish the details of the object, a high resolution image is needed. It is therefore important to pay attention to the computational aspects, for example using the GPU processor as presented in [2]. A more troublesome case is when we consider reconstructing a dynamically changing object, for example a beating heart, or a component on a moving conveyor. This lead to the problem of reconstruction with less measuring time, hence less number of projections.

The Radon Transform (RT), with details introduced in [1], is one of the most commonly used forward modelling

approach that treats the X ray CT projections, with the expression:

$$g(r, \phi) = \mathcal{R}f(x, y) = \int f(x, y) \delta(x \cos \phi + y \sin \phi - r) dx dy \quad (2)$$

where x, y define the position of pixel value in the image, while r, ϕ define the perpendicular length from center point and the angle of X ray under consideration. Based on this transformation and its analytical inversion, there have been many analytical reconstruction methods. We may mention here one of the main methods called Filtered Back Projection (FBP) which can be summarized as $\widehat{f} = \mathcal{B}F^{-1}|\Omega|Fg$ where F and F^{-1} are direct and inverse Fourier transform, $|\Omega|$ is a modulated filter and $\mathcal{B}$ is the Back-Projection operator. In order to solve the reconstruction problem algebraically, the discretization of the model gives the projection rule for each ray i through the pixel j :

$$g[i] = \sum_j H[i, j]f[j] \quad (3)$$

By adding up all the projection rays and all the direction of projections, the synthetic forward operator is available, with the object represented by a vector $\mathbf{f} \in \mathbb{R}^{N \times 1}$, projections represented by vector $\mathbf{g} \in \mathbb{R}^{M \times 1}$ and the linear projection system represented by matrix $\mathbf{H} \in \mathbb{R}^{M \times N}$. Accounting for additive noise of the linear system $\boldsymbol{\epsilon} \in \mathbb{R}^{M \times 1}$, the forward model is:

$$\mathbf{g} = \mathbf{H}\mathbf{f} + \boldsymbol{\epsilon} \quad (4)$$

According to this model, two basic analytical reconstruction methods find their discretized version: Back-Projection (BP) $\widehat{\mathbf{f}} = \mathbf{H}'\mathbf{g}$ and the Filtered Back-Projection (FBP) $\widehat{\mathbf{f}} = \mathbf{H}(\mathbf{H}\mathbf{H}')^{-1}\mathbf{g}$ with analytical and algebraical details presented in [3]. The BP gives a rough blurred reconstruction image, while FBP is an efficient way of reconstruction when there are a great number of projections uniformly distributed around the object. In the cases with strong noise or with insufficient projections, it is no longer robust.

Deterministic Regularization based methods

Many image processing problems are ill-posed. It is necessary to regularize the solution by imposing an a priori constraint. Mathematically, this constraint is often expressed through a regularization function:

$$\mathcal{J}(\mathbf{f}) = \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 + \lambda\mathcal{R}(\mathbf{f}) \quad (5)$$

where λ is called regularization parameter and it controls the tradeoff between data-model adequation and the prior knowledge expressed through the regularization term $\mathcal{R}(\mathbf{f})$. Different choices on the regularization terms are for example:

- L2 (Ridge regression): $\mathcal{R}(\mathbf{f}) = \|\mathbf{f}\|_2^2$
- L2 (Quadratic Regularization): $\mathcal{R}(\mathbf{f}) = \|\mathbf{D}\mathbf{f}\|_2^2$ where $\mathbf{D}$ is a linear operator.
- L1 (Lasso): $\mathcal{R}(\mathbf{f}) = \|\mathbf{f}\|_1$
- Total Variation: $\mathcal{R}(\mathbf{f}) = \|\mathbf{G}\mathbf{f}\|_1$ where $\mathbf{G}$ represents the gradient operator.
- Total Variation for images: $\mathcal{R}(\mathbf{f}) = \|\mathbf{G}_x\mathbf{f}\|_1 + \|\mathbf{G}_y\mathbf{f}\|_1$ where $\mathbf{G}_x$ and $\mathbf{G}_y$ represent correspondingly the horizontal gradient and vertical gradient.

From another point of view, regularization can also impose constraints on the coefficients $\mathbf{z}$ of the signal $\mathbf{f}$ in a proper basis or frame; i.e., $\mathbf{f} = \mathbf{D}\mathbf{z}$, known as synthesis prior presented in [4]:

$$\mathcal{J}(\mathbf{z}) = \|\mathbf{g} - \mathbf{H}\mathbf{D}\mathbf{z}\|^2 + \lambda\mathcal{R}(\mathbf{z}) \quad \text{and} \quad \widehat{\mathbf{f}} = \mathbf{D}\widehat{\mathbf{z}} \quad (6)$$

where $\mathbf{D}$ is the synthesis operator of a decomposition of $\mathbf{f}$ on a dictionary (e.g., gradient, wavelets, etc).

Classical regularization techniques for variational image restoration include constrained Least Square, Tikhonov-Miller method. Since the invention of Total Variation (TV) based regularization, variational image deblurring and its

extensions have received much more attention. The other class of competing schemes in the past decade were sparsity-based regularization. The class of l_1 regularized optimization problems has received much attention because of the presence of an l_1 regularization term, optimization problems are still very difficult to solve. Various of methods have been studied for example the Newton's method presented in [5] and the Split Bregman method presented in [6]. By using these kinds of classical optimization methods, we can get reconstruction results by setting suitable regularization parameters.

Choosing suitable regularization parameter values is necessary and challenging, and very often costly. The ways of choosing this parameter are for example Cross Validation [7] and L-Curve [8] regularization. By adding different regularization terms, l_2 or l_1 , different constraints are considered. l_2 norm criterion gives globally smooth results, and enforce a roughness penalty on the solution. l_1 regularization, on the other hand, enforces the sparsity and therefore preserves edges of image during reconstruction.

Sparsity enforcing prior models

Normally the sparse property of a statistical variable is enforced by using three kinds of distributions:

- The Generalized Gaussian distributions: $p(x) = \frac{\beta}{2\alpha\Gamma(1/\beta)} \exp\left\{-\left(\frac{|x-\mu|}{\alpha}\right)^\beta\right\}$
- The Gaussian Mixture distributions: $p(x) = \sum_j^S w_j \frac{1}{\sqrt{2\pi}\sigma_j} \exp\left\{-\frac{1}{2}\left(\frac{x-\mu_j}{\sigma_j}\right)^2\right\}$
- The heavy tailed distributions, of which the tails are not exponentially bounded.

The standard Student-t distribution is a heavy tailed one, with the expression:

$$St(x|\nu) = \frac{\Gamma\left(\frac{\nu+1}{2}\right)}{\sqrt{\nu\pi}\Gamma\left(\frac{\nu}{2}\right)} \left(1 + \frac{x^2}{\nu}\right)^{-\frac{\nu+1}{2}} \quad (7)$$

From the definition of its variance we easily figure out that there is a limit of the variance of Student-t distribution: $\text{Var}[f] = \frac{\nu}{\nu-2} = 1 + \frac{2}{\nu-2} > 1$. This limit lead to the consequence that this heavy-tailed distribution can't have a small variance, therefore the sparsity couldn't be intensively enforced. In this paper we use a generalization of Student-t distribution (Stg) which can be obtained by using the Normal-Inverse Gamma marginalization property:

$$St_g(f|\alpha, \beta) = \int \mathcal{N}(f|0, v) \mathcal{IG}(v|\alpha, \beta) dv = \frac{\Gamma(\alpha + 1/2)}{\sqrt{2\pi\beta}\Gamma(\alpha)} \left(1 + \frac{f^2}{2\beta}\right)^{-(\alpha + \frac{1}{2})} \quad (8)$$

This generalization of Student-t distribution adds a supplementary parameter to the standard one, and hence is more able to control the sparsity rate of prior distribution.

Linear sparse transformation

In many applications, in particular in Non Destructive Testing (NDT), the objects are piece-wise homogeneous. Noticing that in the homogeneous zone, the gradient is zero, and the non-zero values only exist at the edges of the image. This gives a very convenient property to use the sparsity constraints. A piece-wise continuous image can then be represented by different type of sparse transforms. We use matrix $\mathbf{D}$ to represent the chosen transformation operator which gives a linear relationship between image and the sparse transform coefficients $\mathbf{z}$ with the relationship $\mathbf{f} = \mathbf{D}\mathbf{z}$. In this paper, the multilevel Haar Transformation, which is a form of the Discrete Wavelet Transform (DWT), is used, with benefits that it is an orthonormal operator with the property $\mathbf{D}'\mathbf{D} = \mathbf{I}$ which simplifies the mathematical computations and above all, the inverse operator $\mathbf{D}^{-1}$ and the transpose operator $\mathbf{D}'$ are interchangeable. The multilevel Haar Transformation with level l of an image $\mathbf{f} \in \mathbb{R}^{m \times n}$ is defined as:

$$\text{for } k = 1 \text{ to } l, \quad \mathbf{f}_{\mathcal{H}}^{(k)} = \mathcal{Haar}\left(\mathbf{f}_{\mathcal{H}}^{(k-1)}\left(1 : \frac{m}{2^{k-1}}, 1 : \frac{n}{2^{k-1}}\right)\right), \quad \mathbf{f}_{\mathcal{H}}^{(0)} = \mathbf{f} \text{ and } \mathbf{z} = \mathbf{f}_{\mathcal{H}}^{(l)} \quad (9)$$

Fig.(1) shows the original Shepp-Logan image, its gradient and its corresponding 5-level Haar transformation coefficients. The gradient of the image represents precisely the contours, and similarly for the high frequency coefficients

in the MH transform. By enforcing the sparsity of either the gradient or the high frequency coefficients of the MH transform, the edges of the image could be preserved. The benefit of the MH transform is that it is inversible and orthonormal.

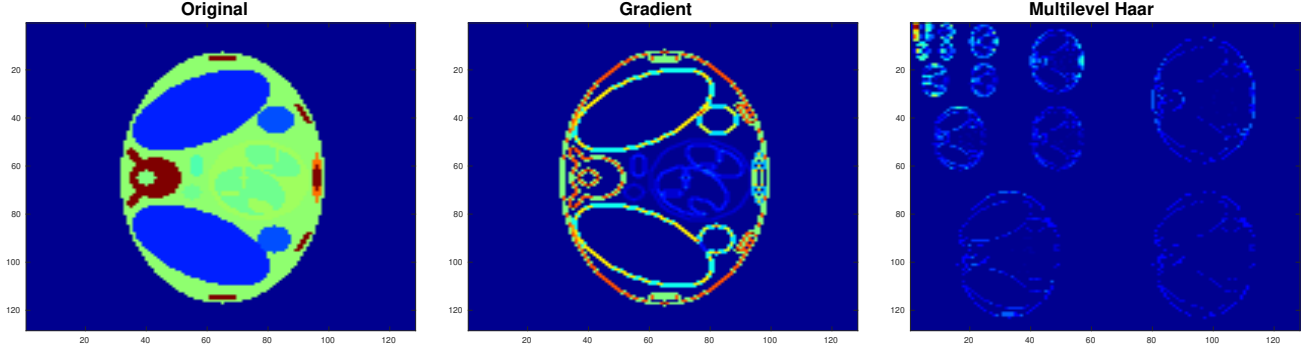


FIGURE 1: Relation of original image, its gradient and its multilevel Haar transform coefficients.

While using this linear dictionary representation, an additive white noise is always considered:

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

The Proposed Unsupervised Hierarchical Haar Transform based Model (HHBM)

The Bayesian methods gives possibilities to do the reconstructions in an unsupervised way, meaning estimating variables as well as parameters. By using the Bayesian method, the prior informations can be combined with the data, and most importantly, it provides a convenient setting for a wide range of models adapting to different prior informations. In this paper, we propose a new hierarchical a priori model for the Bayesian inference in which the sparsity of coefficients of sparse transformation is enforced by using a heavy tailed prior distribution. In this method, both the image and the transformation coefficients are estimated alternatively.

To consider the estimation problem from a Bayesian point of view, we first define a likelihood based on the given data generation model and then we introduce sparsity to our estimate by assigning a suitable heavy-tailed prior distribution over the parameter vector $\mathbf{z}$. Generally, the noise in a linear model is supposed to be a independent and identically distributed (iid) white noise, thus belonging to an i.i.d. Gaussian noise with variance $\mathbf{v}_\epsilon$:

$$p(\boldsymbol{\epsilon}|\mathbf{v}_\epsilon) = \mathcal{N}(\boldsymbol{\epsilon}|\mathbf{0}, \mathbf{V}_\epsilon) \quad \text{where } \mathbf{V}_\epsilon = \text{diag}[\mathbf{v}_\epsilon] \quad (11)$$

According to the forward model shown in Eq.(4), the distribution of the additive noise gives exactly the distribution of the discrepancy of $\mathbf{g}$ and $\mathbf{H}\mathbf{f}$ and hence the likelihood is:

$$p(\mathbf{g}|\mathbf{f}, \mathbf{v}_\epsilon) = \mathcal{N}(\mathbf{g}|\mathbf{H}\mathbf{f}, \mathbf{V}_\epsilon) \quad (12)$$

For the representation of the prior statistical model of image, the multilevel Haar Transformation is used as shown in Eq.(10), where the prior distribution is defined as a Normal distribution while the noise belonging to a Gaussian distribution $p(\boldsymbol{\xi}|\mathbf{v}_\xi) = \mathcal{N}(\boldsymbol{\xi}|\mathbf{0}, \mathbf{V}_\xi)$:

$$p(\mathbf{f}|\mathbf{z}, \mathbf{v}_\xi) = \mathcal{N}(\mathbf{f}|\mathbf{D}\mathbf{z}, \mathbf{V}_\xi) \quad \text{where } \mathbf{V}_\xi = \text{diag}[\mathbf{v}_\xi] \quad (13)$$

The Multilevel Haar Transformation of the image is considered as a sparse representation and a Generalized Student-t (Stg) distribution is used to enforce the sparsity property. The Stg distribution for coefficients $\mathbf{z}$: $p(\mathbf{z}) = \text{Stg}(\mathbf{z}|\alpha_{z_0}, \beta_{z_0})$ can be defined by a hierarchical structure as below:

$$p(\mathbf{z}|\mathbf{0}, \mathbf{v}_z) = \mathcal{N}(\mathbf{z}|\mathbf{0}, \mathbf{v}_z) \quad (14)$$

$$p(\mathbf{v}_z|\alpha_{z_0}, \beta_{z_0}) = \text{IG}(\mathbf{v}_z|\alpha_{z_0}, \beta_{z_0}) \quad (15)$$

where $\alpha_{z_0}, \beta_{z_0} > 0$ are the two hyper-parameters of the hierarchical structured sparsity enforcing Generalised Student-t distribution. In particular, by initializing different α_{z_0} and β_{z_0} value, the scale of sparsity rate of distribution can be controlled.

Lastly, considering the positive property of the variances of noise, as well as that major part of the variance are near zero, the Inverse Gamma (IG) distribution is used to define variables v_ϵ and v_ξ . In this case, the additive noises are therefore belonging to the Generalized Student-t distribution.

A directed acyclic graph (DAG) of the proposed hierarchical Bayesian Haar Transform based model is shown in Fig.(2). The corresponding statistical structure model is on the right.

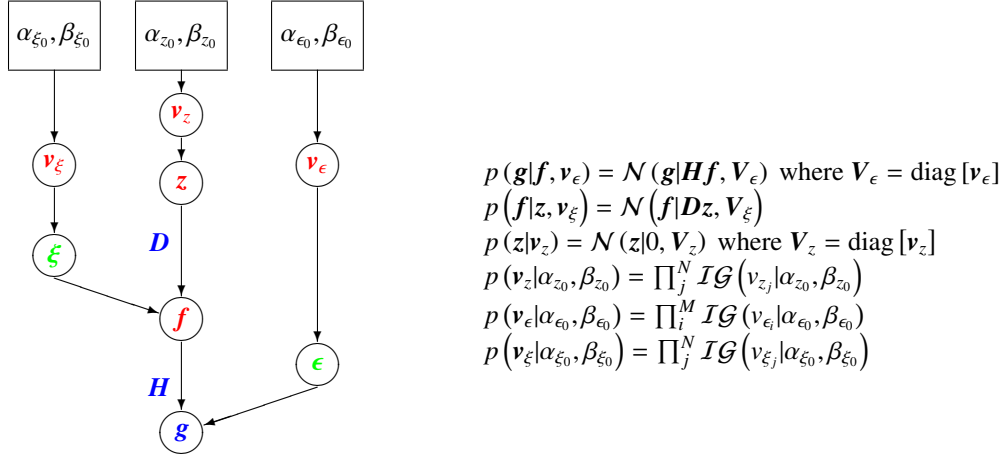


FIGURE 2: Directed acyclic graph of proposed model.

The posterior model and JMAP estimation

So far we have presented the hierarchical structured prior treating the variables and parameters. To proceed with Bayesian inference, the computation of the joint posterior distribution over the variables and parameters is required. Using Bayes' law, this distribution is expressed as shown in Eq.(16):

$$p(\mathbf{f}, \mathbf{z}, \mathbf{v}_z, \mathbf{v}_\epsilon, \mathbf{v}_\xi|\mathbf{g}) \propto p(\mathbf{g}|\mathbf{f}, \mathbf{v}_\epsilon) p(\mathbf{f}|\mathbf{z}, \mathbf{v}_\xi) p(\mathbf{z}|\mathbf{v}_z) p(\mathbf{v}_z|\alpha_{z_0}, \beta_{z_0}) p(\mathbf{v}_\epsilon|\alpha_{\epsilon_0}, \beta_{\epsilon_0}) p(\mathbf{v}_\xi|\alpha_{\xi_0}, \beta_{\xi_0}) \quad (16)$$

Where $\mathbf{z}$ represents the wavelet coefficients of the object $\mathbf{f}$, $\mathbf{D}$ the corresponding transform dictionary, $\mathbf{v}_z$ variance of $\mathbf{z}$, $\mathbf{v}_\epsilon$ variance of noise ϵ and $\mathbf{v}_\xi$ variance of noise ξ .

In this paper, the Joint Maximum A Posterior (JMAP) estimator is applied to estimate the variables and parameters. By using JMAP optimization, all the variables and parameters are optimized alternatively by maximizing the posterior distribution. The optimizing rule is expressed in Eq.(17):

$$[\widehat{\mathbf{f}}, \widehat{\mathbf{z}}, \widehat{\mathbf{v}}_z, \widehat{\mathbf{v}}_\epsilon, \widehat{\mathbf{v}}_\xi] = \arg \max_{\mathbf{f}, \mathbf{z}, \mathbf{v}_z, \mathbf{v}_\epsilon, \mathbf{v}_\xi} p(\mathbf{f}, \mathbf{z}, \mathbf{v}_z, \mathbf{v}_\epsilon, \mathbf{v}_\xi|\mathbf{g}) \quad (17)$$

By estimating alternatively, the updating rule of all the variables and parameters are listed in Eq.(18)-Eq.(22):

$$\text{iter} : \widehat{\mathbf{f}}^{(k+1)} = \widehat{\mathbf{f}}^{(k)} - \widehat{\gamma}_f^{(k)} \nabla \mathcal{J}(\widehat{\mathbf{f}}^{(k)}) \quad (18)$$

$$\text{iter} : \widehat{\mathbf{z}}^{(k+1)} = \widehat{\mathbf{z}}^{(k)} - \widehat{\gamma}_z^{(k)} \nabla \mathcal{J}(\widehat{\mathbf{z}}^{(k)}) \quad (19)$$

$$\widehat{v}_{z_j} = \frac{\beta_{z_0} + \frac{1}{2} \widehat{z}_j^2}{\alpha_{z_0} + 3/2} \quad (20)$$

$$\widehat{v}_{\epsilon_i} = \frac{\beta_{\epsilon_0} + \frac{1}{2} (g_i - [\mathbf{H}\widehat{\mathbf{f}}]_i)^2}{\alpha_{\epsilon_0} + 3/2} \quad (21)$$

$$\widehat{v}_{\xi_j} = \frac{\beta_{\xi_0} + \frac{1}{2} (\widehat{f}_j - [\mathbf{D}\widehat{\mathbf{z}}]_j)^2}{\alpha_{\xi_0} + 3/2} \quad (22)$$

where

$$\mathcal{J}(\mathbf{f}) = \frac{1}{2} (\mathbf{g} - \mathbf{H}\mathbf{D}\mathbf{z})' \mathbf{V}_\epsilon^{-1} (\mathbf{g} - \mathbf{H}\mathbf{D}\mathbf{z}) + \frac{1}{2} (\mathbf{f} - \mathbf{D}\mathbf{z})' \mathbf{V}_\xi^{-1} (\mathbf{f} - \mathbf{D}\mathbf{z}) \quad (23)$$

$$\mathcal{J}(\mathbf{z}) = \frac{1}{2} (\mathbf{f} - \mathbf{D}\mathbf{z})' \mathbf{V}_\xi^{-1} (\mathbf{f} - \mathbf{D}\mathbf{z}) + \frac{1}{2} \mathbf{z}' \mathbf{V}_z^{-1} \mathbf{z} \quad (24)$$

$$\widehat{\gamma}_f^{(k)} = \frac{\|\nabla \mathcal{J}(\widehat{\mathbf{f}}^{(k)})\|^2}{\|\widehat{\mathbf{Y}}_\epsilon \mathbf{H} \nabla \mathcal{J}(\widehat{\mathbf{f}}^{(k)})\|^2 + \|\widehat{\mathbf{Y}}_\xi \nabla \mathcal{J}(\widehat{\mathbf{f}}^{(k)})\|^2} \quad \text{where } \widehat{\mathbf{Y}}_\epsilon = \widehat{\mathbf{V}}_\epsilon^{-\frac{1}{2}} \text{ and } \widehat{\mathbf{Y}}_\xi = \widehat{\mathbf{V}}_\xi^{-\frac{1}{2}} \quad (25)$$

$$\widehat{\gamma}_z^{(k)} = \frac{\|\nabla \mathcal{J}(\widehat{\mathbf{z}}^{(k)})\|^2}{\|\widehat{\mathbf{Y}}_\xi \mathbf{D} \nabla \mathcal{J}(\widehat{\mathbf{z}}^{(k)})\|^2 + \|\widehat{\mathbf{Y}}_z \nabla \mathcal{J}(\widehat{\mathbf{z}}^{(k)})\|^2} \quad \text{where } \widehat{\mathbf{Y}}_z = \widehat{\mathbf{V}}_z^{-\frac{1}{2}} \quad (26)$$

and $\nabla \mathcal{J}(\cdot)$ is the gradient of $\mathcal{J}(\cdot)$.

Computational aspects and GPU implementation

For the updating of $\mathbf{f}$ and $\mathbf{z}$, considering the big data size constraints and the impossibility of calculating the inverse of projection matrix $\mathbf{H}$, the gradient descent optimization algorithm is used, in which the parameters γ_f and γ_z in Eq.(25) and Eq.(26) are the corresponding descent step length. In this paper, considering orthonormal property, $\mathbf{D}' = \mathbf{D}^{-1}$, the matrix $\mathbf{D}'$ corresponds to the multilevel Haar transformation operator, and $\mathbf{D}$ corresponds to the inverse multilevel Haar transformation. The covariance matrix $\mathbf{V}_z$, $\mathbf{V}_\epsilon$ and $\mathbf{V}_\xi$ are diagonal matrix, so that the inverse of them are simply inverse of all the diagonal elements. By using this method, all the calculations can be applied to a 3D CT reconstruction case. The projection, back-projection and Haar transformation operators could also be accelerated by using GPU processor.

Simulation Results

In the simulations, the middle slice image of the Shepp Logan object is used as the original image. The simulation image has size 128×128 and it consists of several homogeneous zones, each of which corresponds to a material. Thus it has a piece-wise continuous property. Parallel projections are applied in angles uniformly distributed from 0 to 180 degree. The reconstruction performance is measured in terms of the nomalized mean square error (NMSE), or the relative error δ_f , which is defined as:

$$\delta_f = \text{NMSE} = \frac{\|\widehat{\mathbf{f}} - \mathbf{f}\|^2}{\|\mathbf{f}\|^2} \quad (27)$$

In Fig.(3), zoon of the reconstructed images by using different methods are listed, with the NMSE of reconstructed images shown below the corresponding image.

By using the Haar Transform base method, while enforcing the sparsity of the sparse transformation coefficients, the edges of the image are better preserved.

In Fig.(4) the profiles of the reconstructed images are compared. The HHBM method outperforms the other method on protecting edges.

The convergence of the relative error of reconstruction of different methods are shown in Fig.(5). Fig.(5)(a) compares different methos by using 128 projections, and Fig.(5)(b) compares the results by using 64 projections. Results proved that the HHBM method stay robust while using insufficient data.

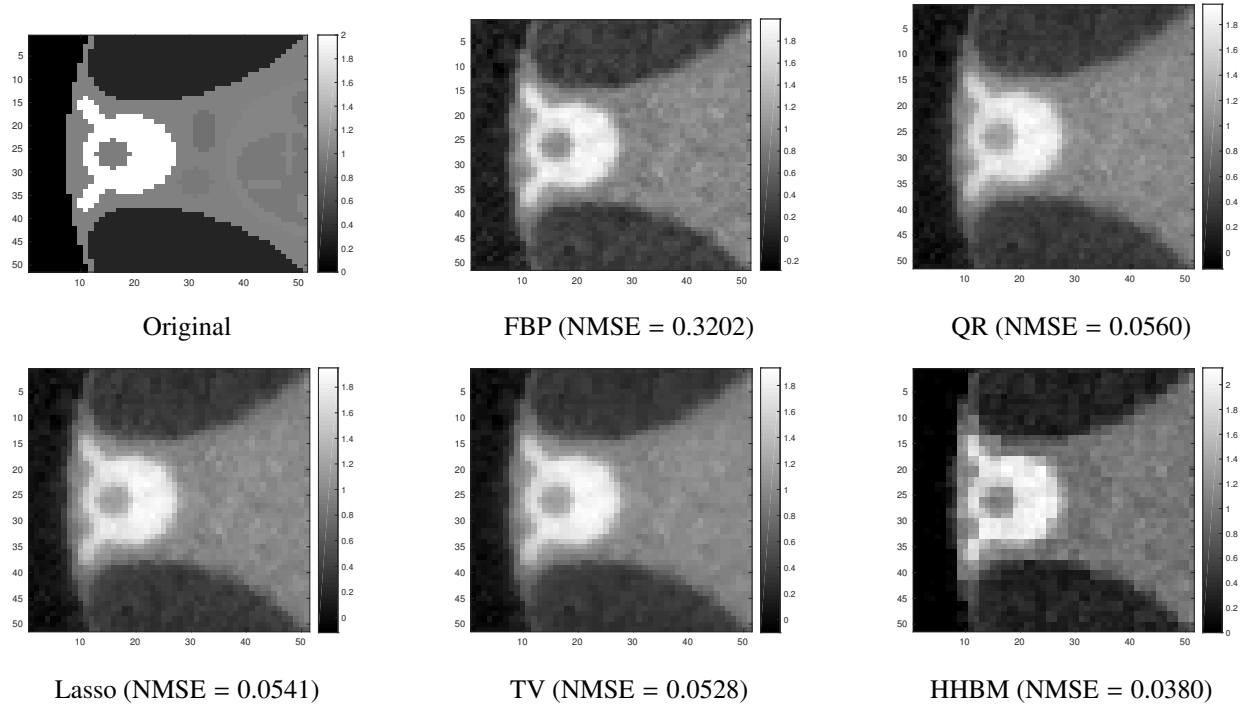


FIGURE 3: Zoon of reconstructed images with different methods: Filtered Back-Projection (FBP), Quadratic Regularization (QR), Lasso, Total Variation (TV) and the proposed HHBM, using 128 projection data and SNR=20dB.

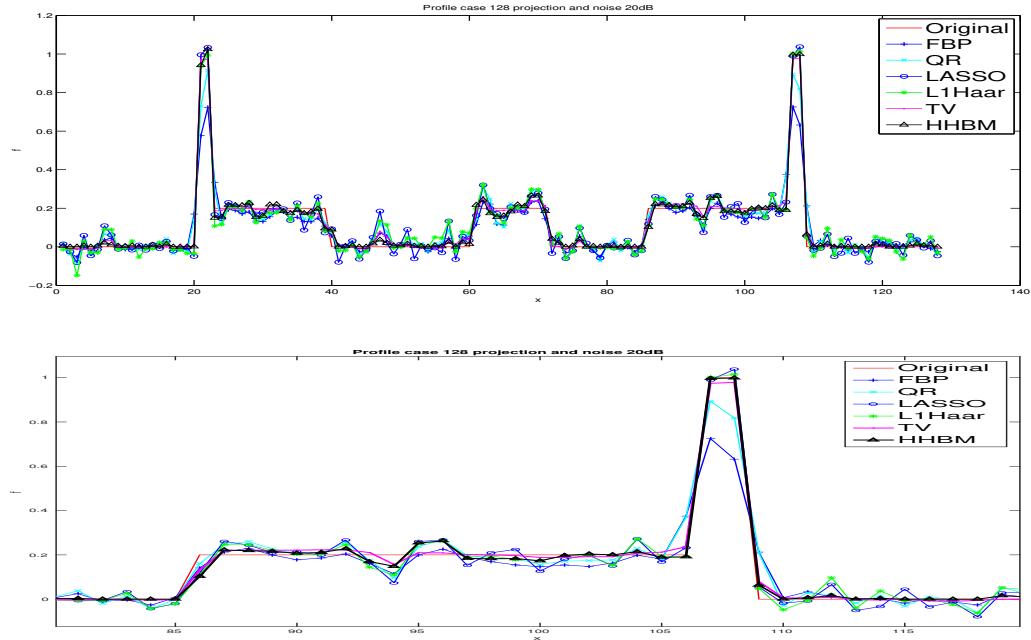


FIGURE 4: Comparison of a) profiles of reconstructed images and b) zone of the profiles. Results obtained within 128 projections and SNR=20dB.

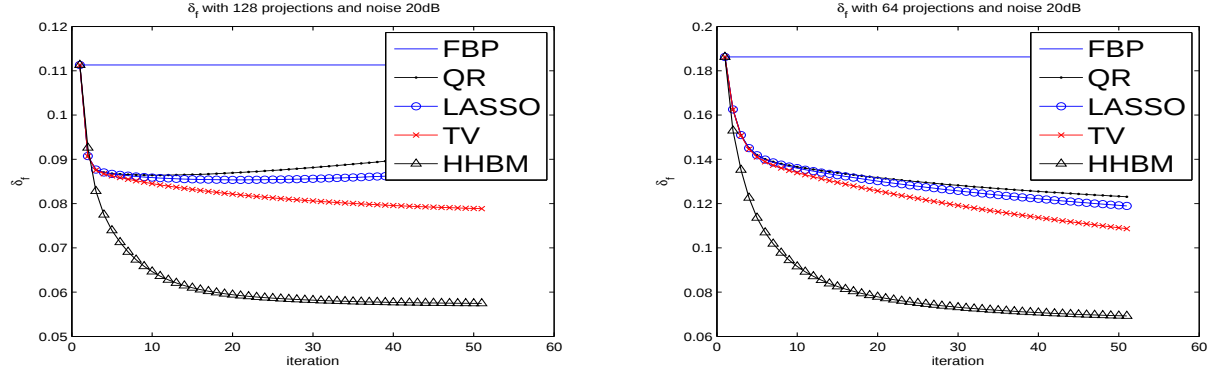


FIGURE 5: Relative error of reconstruction with 20dB noised data and a) 128 projections and b) 64 projections.

Conclusions and Perspectives

In this paper, we have introduced a new hierarchical Haar Transformation based model for the reconstruction of piece-wise continuous X ray Computed Tomography images. Reconstruction using this model enforces the sparsity of the sparse transformation coefficients and hence preserves the edges of reconstructed images. We show that by using our proposed method, the reconstruction of piece-wise continuous image will be obtained with better preserved edges and less projections.

Currently, we are working on the extension of the method to 3D applications, and using a Dual-Tree Discrete Wavelet Transformation as the sparse coefficients. Improved results with respect to noise are expected, due to the robust property of the DT-DWT. Further, we are planing more extended real data tests, including the huge data size industrial Non Destructive Testing one.

REFERENCES

- [1] Rolf Clackdoyle and Michel Defrise. Tomographic reconstruction in the 21st century. *IEEE Signal Processing Magazine*, 27(4):60–80, 2010.
- [2] Nicolas Gac, Alexandre Vabre, Ali Mohammad-Djafari, et al. Multi gpu parallelization of 3d bayesian ct algorithm and its application on real foam reonconstruction with incomplete data set. *Proceedings FVR 2011*, pages 35–38, 2011.
- [3] Ali Mohamad-Djafari. *Inverse problems in vision and 3D tomography*. John Wiley & Sons, 2013.
- [4] Michael Elad, Peyman Milanfar, and Ron Rubinstein. Analysis versus synthesis in signal priors. *Inverse problems*, 23(3):947, 2007.
- [5] Tony F Chan, Gene H Golub, and Pep Mulet. A nonlinear primal-dual method for total variation-based image restoration. *SIAM journal on scientific computing*, 20(6):1964–1977, 1999.
- [6] Tom Goldstein and Stanley Osher. The split bregman method for l1-regularized problems. *SIAM journal on imaging sciences*, 2(2):323–343, 2009.
- [7] Ron Kohavi et al. A study of cross-validation and bootstrap for accuracy estimation and model selection. In *Ijcai*, volume 14, pages 1137–1145, 1995.
- [8] Per Christian Hansen and Dianne Prost O’Leary. The use of the l-curve in the regularization of discrete ill-posed problems. *SIAM Journal on Scientific Computing*, 14(6):1487–1503, 1993.