

Received November 1, 2019, accepted December 8, 2019, date of publication December 12, 2019, date of current version December 23, 2019.

Digital Object Identifier 10.1109/ACCESS.2019.2959037

A Deep Learning Framework for Transforming Image Reconstruction Into Pixel Classification

KAMLESH PAWAR^{1,2}, ZHAOLIN CHEN¹, (Member, IEEE), N. JON SHAH^{1,3}, AND GARY. F. EGAN^{1,2}

¹Monash Biomedical Imaging, Monash University, Clayton, VIC 3800, Australia

²School of Psychological Sciences, Monash University, Clayton, VIC 3800, Australia

³Research Centre Jülich, Institute of Medicine, 52425 Jülich, Germany

Corresponding author: Kamlesh Pawar (kamlesh.pawar@monash.edu)

ABSTRACT A deep learning framework is presented that transforms the task of MR image reconstruction from randomly undersampled k-space data into pixel classification. A DL network was trained to remove incoherent undersampling artifacts from MR images. The underlying, fully sampled, target image was represented as a discrete quantized image. The quantization step enables the design of a convolutional neural network (CNN) that can classify each pixel in the input image to a discrete quantized level. The reconstructed image quality of the proposed DL classification model was compared with conventional compressed sensing (CS) and a DL regression model. The reconstructed images using the DL classification model outperformed the state-of-the-art compressed sensing and DL regression models with a similar number of parameters assessed using quantitative measures. The experiments reveal that the proposed deep learning method is robust to noise and is able to reconstruct high-quality images in low SNR scenarios where conventional CS reconstructions and DL regression networks perform poorly. A generic design framework for transforming MR image reconstruction into pixel classification is developed. The proposed method can be easily incorporated into other DL-based image reconstruction methods.

INDEX TERMS Magnetic resonance imaging, compressive sensing, deep learning.

I. INTRODUCTION

Accelerating data acquisition in magnetic resonance imaging has been an active area of research since the inception of MRI. Initial efforts were made to accelerate the data acquisition process by developing fast imaging sequences, such as echo planar imaging [2] and gradient echo sequences [3]. However, further acceleration of MR data acquisition using these fast imaging sequences resulted in nerve stimulation due to the rapid switching of the magnetic field gradients. Multiple receive channels, which were initially intended to improve the SNR of MRI, were later found to have more valuable application in accelerating MR data acquisition. Accelerated MR acquisitions use multiple receiver channels with coils placed at spatially different locations around the patient. The underlying MR signals are modulated as a result of the respective coil sensitivities. Parallel imaging techniques [4]–[8] exploit knowledge of the variable coil sensitivities present in multichannel receiver coils to reconstruct

MR images from accelerated, undersampled k-space data. Compressed sensing (CS) [9], [10] is another technique that exploits the signal sparsity inherent in MR images. MR images are known to be sparse in a number of signal domains, including wavelet and total variation (TV). Compressed sensing uses knowledge of this sparsity and reconstructs the images from undersampled k-space data using iterative reconstruction methods. In conventional CS-MRI reconstruction, a universal sparsifying transform is assumed (such as wavelet or TV) that transforms the underlying image into a sparse representation. Given the data consistency constraint, the iterative reconstruction algorithm minimizes a cost function to recover this sparse representation. CS-MRI reconstruction is a computationally-intensive process and efforts have been made to develop faster reconstruction methods [11]–[14] by optimizing the algorithms or implementing them on graphics processing units (GPU).

Deep learning [15], [16] is another rapidly growing area finding utility in numerous imaging applications. Deep learning, or more specifically convolutional neural networks

The associate editor coordinating the review of this manuscript and approving it for publication was Ravibabu Mulaveesala¹.

(CNN) [15], have shown remarkable performance improvements in a multitude of computer vision tasks, such as image classification [1] and image segmentation [16]–[19]. Deep learning methods have even surpassed human-level performance in some image classification tasks [1]. Recently, the application of deep learning to image reconstruction from undersampled k-space MR data has been explored by a number of research groups [20]–[27]. In [21], a deep convolutional neural network for solving an inverse problem was presented with an emphasis on CT image reconstruction. In [22], a deep learning method based on the Unet [18] architecture was presented which learns the undersampling artifact instead of the image to enable these artifacts to be removed from the corrupted image. In [24], a method using a cascade of CNN for image reconstruction from undersampled data was presented, which is effectively an unrolling of the iterative reconstruction process. In [28] a framework based on convolutional framelets was presented, linking the mathematical signal processing concepts to the deep learning reconstructions. In [27], a method of reconstruction that maps the undersampled k-space to image space was presented. However, this method required a large amount of internal memory and implementation of the method was only demonstrated for small image sizes of 128×128 . For larger image sizes, such as 256×256 or 512×512 , the implementation poses significant technical challenges. These deep learning methods are motivated by the fact that iterative reconstruction can be converted to non-iterative CNN inferences.

Recent advances in deep learning are being driven by computer vision research in which the image processing task is usually a classification and/or segmentation task. The output of a deep learning network for an image classification/segmentation task is a probability estimate of the object belonging to a certain class. However, the task of image reconstruction is significantly different from a classification/segmentation task. In image reconstruction, the DL network is trained to learn a nonlinear mapping between the input (images reconstructed from undersampled k-space data) and the output (images reconstructed from fully sampled k-space data). DL networks for image reconstruction are regression networks, whereas DL networks for classification/segmentation are classification networks.

The data acquired during MR data acquisition and reconstructions are represented in floating-point numbers. However, at the end of the imaging process, the images are stored and viewed in a dicom format. In the dicom format, images are quantized and consist of a fixed number of grey levels determined by the number of bits used to store the image (usually 12 to 16 bits). Motivated by the fact that MR images are finally stored as quantized images in this work, we present a method to directly reconstruct the quantized images by transforming the task of image reconstruction into a pixel classification task. Image reconstructions from a deep learning classification network were compared with image reconstructions from compressed sensing and a deep learning regression network. We investigated the effect of noise on

the image reconstruction and quantitatively compared the resultant images for low signal to noise acquisitions. The following are the major contributions of this paper:

- A novel DL framework is introduced to model the image reconstruction problem into a pixel classification problem.
- A divide-and-conquer approach is developed for application of the DL classification network to high-bit precision (i.e. 16 bit).
- The pixel classification approach with compressed sensing and the conventional DL regression approach are validated using three experiments which include T1 and T2 weighted MR images, T1 and T2 images with added noise, and an unseen tumor dataset.

II. BACKGROUND

A. ITERATIVE RECONSTRUCTION - COMPRESSED SENSING MRI

The pulse sequence in MRI encodes the NMR signal into the Fourier space using magnetic field gradients. The acquired signal in k-space is a Fourier transform of the image. The data acquisition can be represented as

$$y = Fx,$$

where $y \in C^{N \times N}$ is the acquired k-space data, $x \in C^{N \times N}$ is the $N \times N$ image and $F \in C^{N \times N}$ is a 2D Fourier transform operator. For Cartesian CS acquisition, the acquired k-space data is undersampled along the phase encoding (PE) direction for 1D undersampling and in both the PE and slice encode (SE) direction for 2D undersampling. The undersampled data can be represented as:

$$y_u = U \odot (Fx), \quad (1)$$

where $\odot$ is an elementwise matrix product and U is an undersampling matrix of size $N \times N$, with a value of 1 at the locations where the k-space is sampled and 0 elsewhere.

Consider an image x that can be sparsely represented in the domain ψ , i.e. $x_s = \psi x$ is sparse. The acquired undersampled data (y_u), in terms of a sparse representation of the underlying image can be written as

$$y_u = U \odot (F\psi^{-1}x_s), \quad (2)$$

where ψ^{-1} is the inverse of ψ .

The CS reconstruction minimizes a cost function to reconstruct an estimate of the fully sampled image $\hat{x} = \psi^{-1}\hat{x}_s$ from undersampled k-space data y_u , given by

$$\min_{\hat{x}_s} \lambda \|\hat{x}_s\|_1 + \|y_u - U \odot (F\psi^{-1}\hat{x}_s)\|_2^2 \quad (3)$$

where $\|\cdot\|_1$ and $\|\cdot\|_2$ are l_1 and l_2 norms, respectively. The first term in the cost function enforces sparsity, whereas the second term enforces the consistency of the estimated $\hat{x}_s$ with the acquired data, y_u . The regularization parameter λ determines the level of sparsity and is usually determined empirically, depending on the noise level in the acquired data.

B. CONVOLUTIONAL NEURAL NETWORK (CNN) RECONSTRUCTION

In this section, we briefly describe the convolutional neural network architecture. The topic of CNN is broad and this section only covers the fundamental principles required for CNN image reconstruction. The basic principle of a CNN [15] is a set of convolution operations followed by the addition of a scalar bias term which is then followed by a nonlinear squashing function. In a typical CNN architecture, the input is first convolved with a 2D kernel (width $\times$ height) of a specified size across all the different channels. A bias is added to the convolution output and this output is passed through a nonlinear activation function, typically a rectilinear unit (ReLU). The operations of convolution and nonlinear activation represent a layer of the CNN architecture. There are a large number of different kernels in each layer of a CNN, with the output of the layers being called feature maps. A CNN architecture for an MR image reconstruction consists of many such layers, with the weights of the convolutional kernels learned from a set of training data (consisting of images from fully sampled and undersampled k-space data) using the backpropagation algorithm.

III. PROPOSED METHOD

A. TRANSFORMING IMAGE RECONSTRUCTION INTO PIXEL CLASSIFICATION

The problem of image reconstruction from undersampled data using deep learning has been demonstrated in [20]–[22], [24], [26], [27]. In the demonstrated methods, the CNN architectures were modeled as a regression model, which learns a nonlinear relationship between an image from undersampled k-space data and an image from fully sampled k-space data.

In this work, we propose a DL method to convert the problem of image reconstruction from undersampled data into a pixel classification problem. In order to convert the image reconstruction problem into pixel classification, the target image is first quantized to a finite n-bit discrete grey-level image (Figure 1 (c)). The quantization step results in 2^n unique pixel intensity values. The discretization of the target image makes it possible to design a classification CNN architecture that can classify each pixel in the image from undersampled k-space data to one of the 2^n discrete grey levels. The classification CNN predicts the probability of each pixel belonging to one of the different 2^n classes. The training of the classification network requires the categorical loss in the backpropagation step and the model learns a nonlinear classification model between pixels in the input image (image with undersampling artifact) and the output image (image without artifact).

Deep learning based MR image reconstruction methods use the process of training to fit a nonlinear model by relating an image with undersampling artifact to an image without the artifact.

$$\hat{x} = G_{\theta}(F^*y_u) \quad (4)$$

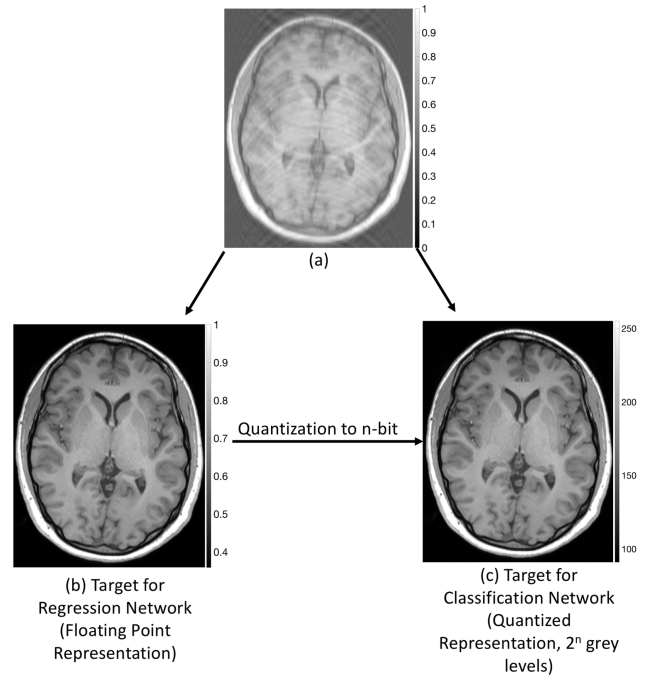


FIGURE 1. The Fundamental principle for transforming image reconstruction into a pixel classification problem. (a): MR image with undersampling artifacts; (b): Target image for regression network represented in 32-bit floating; (c): Target image for classification network represented in an unsigned 8-bit integer.

where, $\hat{x}$ is an estimate of a fully sampled image, y_u is undersampled k-space data, F^* is a 2D inverse Fourier transform operator and G_{θ} is a DL nonlinear model parameterized by θ and generated through the process of training. G_{θ} is an L-layer configuration determined by the architecture of DL-CNN.

Conventionally, the regression DL architectures are used for MR image reconstruction [21], [22], [24], [27] as this minimizes reconstructed image loss such as mean squared error or mean absolute error using training datasets. The mean absolute error for the regression network is defined as

$$\min_{\theta} \mathcal{L}(\theta) = \|x - G_{\theta}^R(F^*y_u)\|_1 \quad (5)$$

In the above equation, G_{θ}^R is a regression DL model, termed DL-R in this paper, and x is the ground truth image from fully sampled k-space data. The loss function aims to minimize the total l_1 error of the reconstructed image. Conventional DL-R image reconstruction models show reduced reconstruction error compared with CS reconstruction. However, residual noise and image blurring are still present in the reconstructed images.

Instead of directly reconstructing an image $\hat{x}$, we propose using the classification DL network for MR image reconstruction, termed DL-C in this paper. The DL-C networks predict the probability of each pixel belonging to the discretized intensity level, $c \in [0, M - 1]$, and M represents the total number of classes. The total number of classes is determined by the bit depth used for quantization ($M = 2^n$).

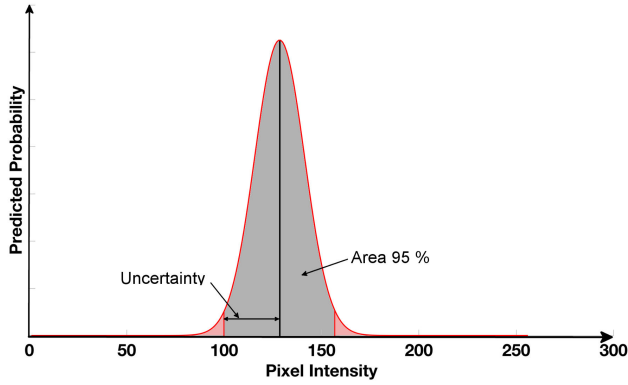


FIGURE 2. A typical predicted probability distribution from the classification network for a given pixel. The x-axis represents quantized pixel intensities and the y-axis represents the predicted probability. The uncertainty is defined as the distance from the mean values at 95% confidence.

The probability of an output pixel, P , falling into the intensity level, c , is given by

$$\rho(c) = G_{\theta}^C(F^*y_u), \quad (6)$$

where G_{θ}^C is the classification DL model. Since the regression problem of image reconstruction is transformed into a pixel classification problem, the loss function to be minimized must be from a categorical loss function. Hence, in this work, we minimized the categorical cross entropy defined as

$$\min_{\theta} \mathcal{L}(\theta) = - \sum_{c=0}^{M-1} \rho_t(c) \log(\rho_p(c)), \quad (7)$$

where, $\rho_t(c)$ is the true probability for the class c , obtained from training datasets and $\rho_p(c)$ is the predicted probability for the class c for the model G_{θ}^C . The categorical cross entropy is a probabilistic loss function and is independent of the individual pixel value. This loss function penalizes all the pixels equally, irrespective of the pixels' values, and the loss only depends on the predicted probability distribution of the pixel.

The adjacent classes or labels (i.e pixel intensities) in the image reconstruction task results in the predicted probability $\rho_p(c)$ distribution being nearly symmetrical and Gaussian-like in nature (Figure 2). The approach is more robust at the sharp edges and for noise with zero means. This is because the presence of tissue boundaries or zero-meaned noise, $\rho_p(c)$, can be expected to have wider distributions while the mean remains largely undisturbed. In other words, the presence of tissue boundaries or noise results in an increase in the uncertainty (Figure 2), but the mean largely remains the same and hence the pixel intensity can be recovered more accurately.

B. Quantization

The quantization [29]–[32] is a step to discretize the signal from a large number of values to a smaller fixed number of values. A signal can be quantized using a uniform or non-uniform discretization process. In this work, we employed a

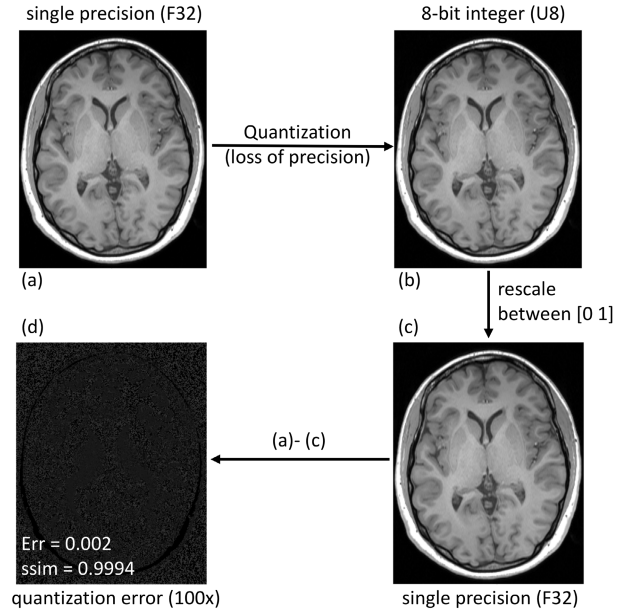


FIGURE 3. Analysis of the quantization error; (a): 32-bit single-precision floating-point image scaled between [0,1]; (b): an unsigned 8-bit integer representation of the image generated from (a), this step introduces quantization error due to fixed point conversion; (c): image (b) converted to single-precision floating-point by scaling (dividing each pixel by 255); (d): quantization error image (scaled by 100x for viewing) generated by subtracting the original image (a) and the quantized version (c). Relative error due to quantization is 0.002 and the structural similarity index between images (a) and (c) is 0.9994.

uniform quantization approach to convert the floating-point representation of an image I_{float} , scaled in the range [0, 1], to a fixed-point representation. The uniform quantization to n -bits can be performed using

$$I_{fixed} = g[(2^n - 1) \cdot I_{float} + 0.5], \quad (8)$$

where $g[\cdot]$ is the ceiling function and I_{fixed} is an n -bit fixed point representation of the image.

The quantization step introduces an error into the image, and the information lost due to quantization cannot be recovered. Therefore, theoretically, the quantization error is the minimum error that is always present in the reconstructed image using the DL classification framework. The quantization error in all of the MR images is empirically quantified by discretizing the floating-point representation to an 8-bit unsigned representation. An example of the quantization process is demonstrated for an MR image in Figure 3. The initial ground truth image is scaled between [0, 1] and represented in a single precision floating point (Figure 3 (a), I_{float}). The image is quantized to an 8-bit fixed point unsigned integer number [0, 255] (Figure 3 (b), I_{fixed}). The quantization step introduces errors in the image, and the quantized image (fixed point) is again converted to a single precision floating point number and scaled between [0, 1] (Figure 3 (c), $\hat{I}_{float}$). Although the quantized image is now represented as a floating point, it only consists of 256 different intensity levels. Figure 3 (d) shows the absolute error introduced due

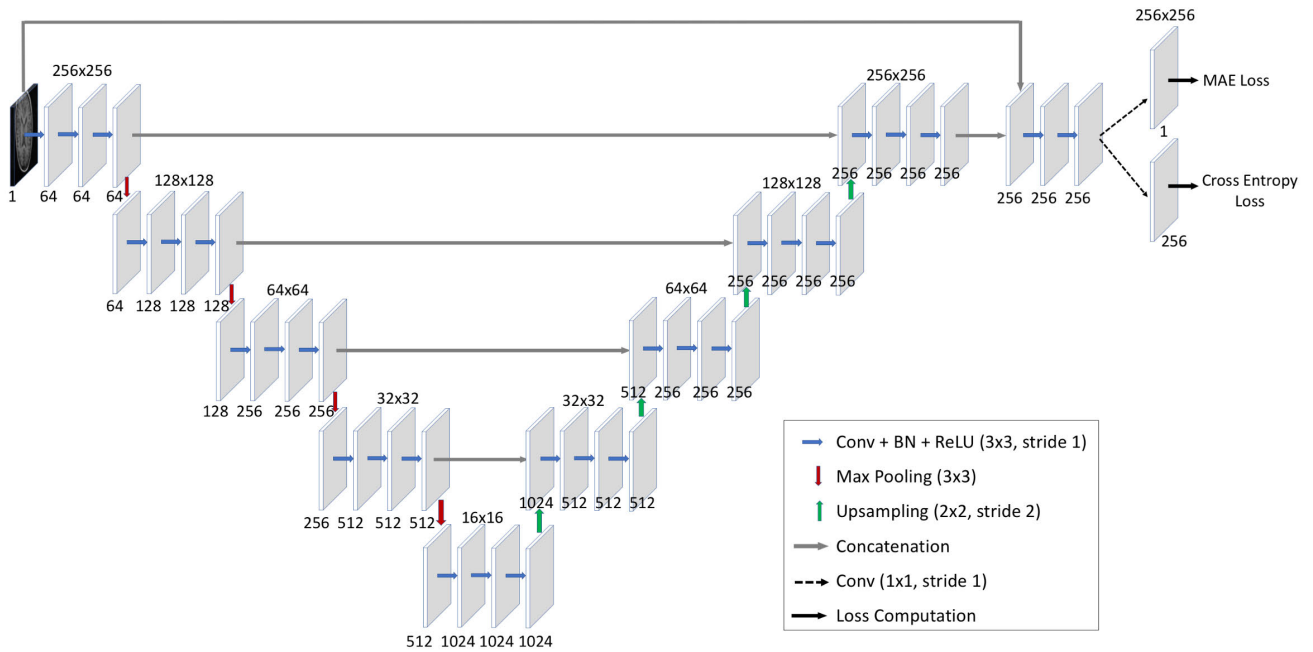


FIGURE 4. The Network Architecture: Residual Encoder-Decoder convolution neural network architecture for image reconstruction. The input to the network is a 256×256 image with undersampling artifact. The numbers at the top of the filters show the size of the data and the numbers at the bottom show the number of filtered outputs. In the case of the classification model, the last layer is a 256-feature layer that predicts the probability of quantized image intensity level for each pixel in the reconstructed image. In the case of the regression model, the last layer is a single feature layer that predicts the floating-point value of each pixel in the reconstructed image. The figure represents the architecture for both the regression network and the classification network - as differentiated by the dashed arrow. Only one of the dashed arrow layers was included to form either a classification or a regression network. Only the top (MAE Loss) layer was included for the regression network, while only the bottom (cross-entropy loss) layer was included for the classification network.

to the process of quantization - here the relative error was 0.002 or 0.2%. This example demonstrates that the error introduced due to quantization is small and does not affect the visual appearance of the image. Theoretically, the quantization error [33] for n -bit quantization can vary from $-\text{LSB}/2$ to $+\text{LSB}/2$ (LSB: least significant bit). The LSB for 8-bit quantization is $1/2^8$, therefore the maximum error per pixel can be $\pm 1/2^9 = \pm 0.002$ or $\pm 0.2\%$.

C. DEEP LEARNING NETWORK ARCHITECTURE

An encoder-decoder Unet [18], [28] architecture (Figure 4) with skip connection was designed and trained to predict an estimated probability distribution for each pixel intensity in the artifact-free image from the image with the undersampling artifacts. The encoder network consists of a series of convolution and pooling layers, with three convolutions before every pooling operation. The decoder consists of a series of convolution and up-sampling layers, with three convolutions before every up-sampling operation. The decoder network consists of dropout layers after every convolution, with a dropout fraction of 0.2. The base architecture for both the regression and classification models remains the same except for the last layer and the loss function.

For the regression network, the last layer uses the mean absolute error as the loss function to predict a target output image. The loss function computes the mean absolute error between the output and the target image (artifact-free

floating-point image). The input to the network is a 256×256 image from the undersampled data and the output is an estimate of the image from fully sampled data.

For the classification network, the last layer consists of the softmax activation and cross entropy as a loss function. The loss function computes the categorical cross-entropy loss between the output probabilities of the network and the target probabilities. The target probabilities are determined by the pixel intensity in the quantized image, i.e. the probability of the class corresponding to the pixel intensity will be one. The input to the network is a 256×256 undersampled floating-point image and the output is the class probabilities for the pixels. The class index of the maximum predicted probability can be assigned as the value of the given pixel. The resulting quantized image is transformed to a floating-point number and scaled between $[0,1]$. Although the final output image from the DL classification network is represented in floating-point, the number of grey levels is discrete and is determined by the bit-depth used for the initial quantization.

D. IMAGE RECONSTRUCTION FROM PREDICTED PROBABILITY

The proposed DL classification method outputs a probability distribution for each pixel instead of the exact pixel value (Figure 2). In this work, we use two reconstruction methods: a) DL-C Max and b) DL-C Avg. In DL-C Max, we employed a simple method of image reconstruction to assign the class

index of the maximum probability as the value to the given pixel at spatial location r . This results in an image with $M = 2^n$ discrete values, determined as:

$$I(r) = \arg \max_c \rho_p(c, r). \quad (9)$$

The DL-C Avg takes the weighted sum of the probability distribution, determined as:

$$I(r) = \sum_{c=0}^{M-1} c \cdot \rho_p(c, r), \quad (10)$$

where $I(r)$ is the reconstructed image and $\rho_p(c, r)$ is the probability distribution predicted by the DL classification network at spatial location r in the image for $M = 2^n$ classes, where n is the number of bits used for quantization. The weighted sum approximates continuous values between $[0, M - 1]$ instead of discrete M values, and thus minimizes the error due to quantization.

E. DIVIDE-AND-CONQUER APPROACH FOR HIGHER BIT PRECISION PIXEL CLASSIFICATION

The memory requirement for the network architecture in section-III-C increases exponentially with the bit-depth. The dynamic range of an n -bit unsigned number N_{n_bit} is $[0, 2^n - 1]$. In order to predict the probability of each possible pixel intensity using the proposed pixel classification, the approach requires 2^n channels at the output of the network. The number of channels can be extremely large, for instance for a 16 bit-depth it is 65536. The large number of channels may not be practical due to physical memory limitation and the model may not generalize due to overfitting from an increased number of parameters.

In order to make the classification network practical for higher bit precision, we used a divide-and-conquer paradigm. The basic idea of a divide-and-conquer paradigm is to split the large problem into two or more smaller sub-problems. The solution to the individual sub-problems can then be combined to arrive at the solution to the larger problem.

We present a novel network architecture based on a divide-and-conquer paradigm to predict the higher bit-depth images without any substantial increase in the number of learnable parameters. In order to reduce the number of parameters required for higher bit-depth prediction, we propose splitting the n -bit integer number, N_{n_bit} , in to a linear combination of two $n/2$ -bit numbers as:

$$N_{n_bit} = (2^{n/2} a_0 + a_1); \quad \forall a_0, a_1 \in [0, 2^{n/2} - 1] \quad (11)$$

For instance, a 16-bit integer number can be represented as:

$$N_{16_bit} = (256 a_0 + a_1); \quad \forall a_0, a_1 \in [0, 255] \quad (12)$$

As described in equation (11), both a_0 and a_1 belong to a $n/2$ -bit number, therefore the network in Fig.4 can be modified to predict both a_0 and a_1 , and equation (11) can be used to compute the n -bit value N_{n_bit} . This approach does not result in any significant increase in the number of parameters and makes the proposed method practical for higher bit images.

For instance, in order to predict a 16-bit image using this approach, the number of channels required in the output layer is only 256 instead of 65536.

Figure 5 shows the network used to predict the 16-bit precision pixel values using the proposed pixel classification approach. It consists of two outputs a_0 and a_1 , and the 16-bit pixel value is calculated using equation [1].

IV. EXPERIMENTAL RESULTS

A. DATA PREPARATION, NETWORK TRAINING AND VALIDATION

T1 and T2 weighted 3D volumetric images from the IXI dataset (<https://brain-development.org/ixi-dataset/>) were used to train the DL models. For each contrast, T1 and T2, we used 260 subjects for training the network and 64 subjects for validation. The training images were generated by undersampling the k-space data by a factor of 8 in two of the encoding directions with a variable density undersampling pattern. Since the data was 3D, it was possible to perform 2D undersampling with a high acceleration factor of 8, (Acceleration factor = (total number of k-space points) / (number of acquired k-space points)), while for 1D undersampling at such an acceleration factor would result in images with higher artifacts. All the 3D volume images were normalized by dividing both the input and output by the maximum value of the input volume, which scales the data between 0 and 1. The Keras deep learning library with the Tensorflow backend, was used for training the networks. The Adam optimizer was used with initial learning rate = 0.0001, and the learning rate was annealed by a factor of 0.96 for each epoch. One epoch consisted of 2500 iterations. A total of 150 epochs were used for training, and the model for which validation loss was minimum was selected as the final trained model.

Four different methods of MR image reconstruction from undersampled k-space data were compared: (i) a compressed sensing reconstruction with a wavelet and total variation (TV) penalty; (ii) a DL regression reconstruction with mean squared loss; (iii) a DL regression reconstruction with mean absolute loss; and (iii) a DL classification reconstruction (16-bit precision).

Three validation experiments were designed to test the performance of the DL models using the IXI datasets, the IXI datasets with added noise, and the brain tumor cases from the BRATS dataset [31]. Details are provided in Section IV-C.

B. IMAGE RECONSTRUCTION

For the CS reconstruction, the wavelet and the TV regularization parameters were empirically optimized to provide maximum SSIM. The regularization parameters for the wavelet and TV penalties were 0.0001 and 0.0005, respectively. A nonlinear conjugate gradient method was used for the CS reconstruction. The last layer directly predicted the pixel value for the DL regression network reconstruction. The last layer predicted the probability of each pixel belonging to one of the 256 different classes for the DL classification

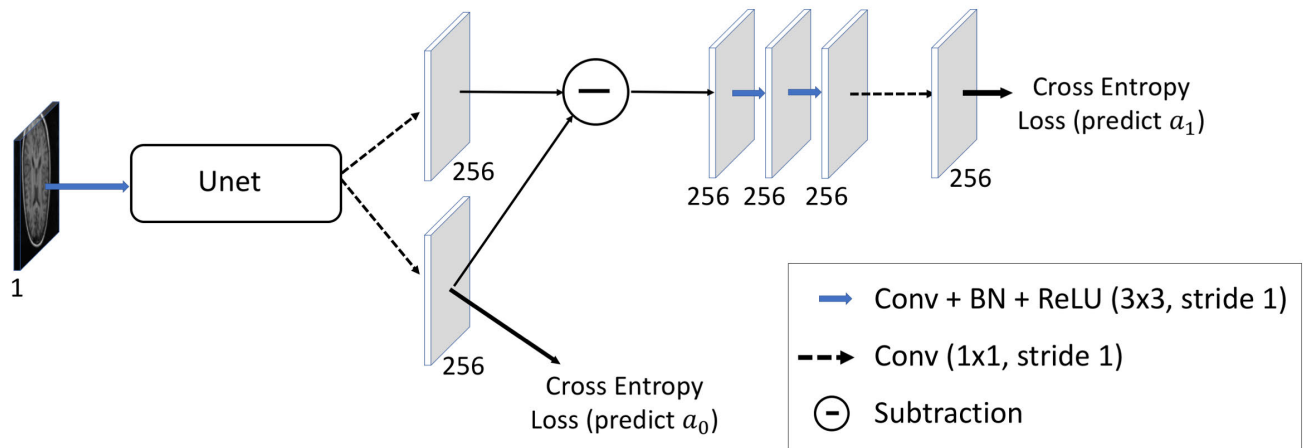


FIGURE 5. The network architecture for 16-bit classification network. The output from the Unet consists of two outputs a_0 and a_1 . Both a_0 and a_1 are 256 channel outputs predicting the probability distribution of two 8-bit numbers. Both the probability distributions are used to compute the final value of a pixel using equation 12.

network. The estimated values for coefficients a_0 and a_1 were calculated using equation (10). Once the coefficients a_0 and a_1 were known, the final 16-bit images were calculated using equation 12.

We further evaluated the performance of the three reconstruction methods in the presence of noise. Random Gaussian noise was added to the undersampled k-space data and resulted in Rician distributed noise in the magnitude images. The zero-filled noise corrupted MR images were provided as input to the three reconstruction methods and the results were compared. The hyperparameters for the CS reconstruction were separately optimized for the noise case.

C. RESULTS

A total of five different reconstructions were compared for the acceleration factor of 8. These were ZF: zero-filled reconstruction where missing k-space lines were simply filled with zeros and the image is reconstructed with the Fourier transform of k-space, CS: compressed sensing reconstruction with TV and wavelet regularizations, DLR-L2: deep learning regression reconstruction with the network trained with mean squared loss, DLR-L1: deep learning regression reconstruction with the network trained with mean absolute loss and DLC: deep learning classification reconstruction (16-bit).

Separate networks were trained for the T1 and T2 weighted images. On visual inspection of Figure 6, the compressed sensing image reconstructions shows a loss of resolution and suffers from blurred edges (enlarged images sections in Figure 6). The DLR-L2 reconstruction resulted in loss of resolution and blurry edges, while the DLR-L1 image reconstructions performed poorly in recovering low contrast features in the reconstructed image. The low contrast features were washed off by the DLR network. The small structure, as pointed out by yellow arrow (Figure 6), shows that the delineation of the white matter and gray matter is more prominent

TABLE 1. Quantitative scores for T1 weighted images.

Methods	SSIM	PSNR	Relative Error	MSE
Compressed Sensing	0.9419	30.66	0.0555	8.583e-4
DLR-L2	0.9484	31.74	0.0491	6.690e-4
DLR-L1	0.9519	31.89	0.0482	6.465e-4
DL Classification	0.9557	32.19	0.0466	6.034e-4

TABLE 2. Quantitative scores for T2 weighted images.

Methods	SSIM	PSNR	Relative Error	MSE
Compressed Sensing	0.9587	31.26	0.0530	7.466e-4
DLR-L2	0.9659	32.71	0.0449	5.361e-4
DLR-L1	0.9678	33.33	0.0418	4.664e-4
DL Classification	0.9723	33.46	0.0410	4.446e-4

in the DLC reconstructed images. The proposed DLC network showed higher resolution, sharper edge images and was able to faithfully recover low contrast features compared to the other reconstructions methods. The superior performance of the DLC network was also evident from the quantitative score of structural similarity (SSIM) index, peak signal to noise ratio (PSNR), relative error and mean squared error (MSE) as demonstrated in Tables 1 and 2.

At low field strengths, the SNR is lower than the high field strength. In order to test the performance of all three algorithms in a low SNR scenario, we trained both the DLC and DLR networks on the dataset with the addition of Gaussian noise. The Gaussian noise was complex valued (both real and imaginary parts drawn from Gaussian distribution) and was added to the undersampled k-space data. The peak signal to noise ratio after adding the noise was approximately 18.0 dB. The presence of noise severely affected the performance of the CS reconstruction and resulted in excessive blurring and

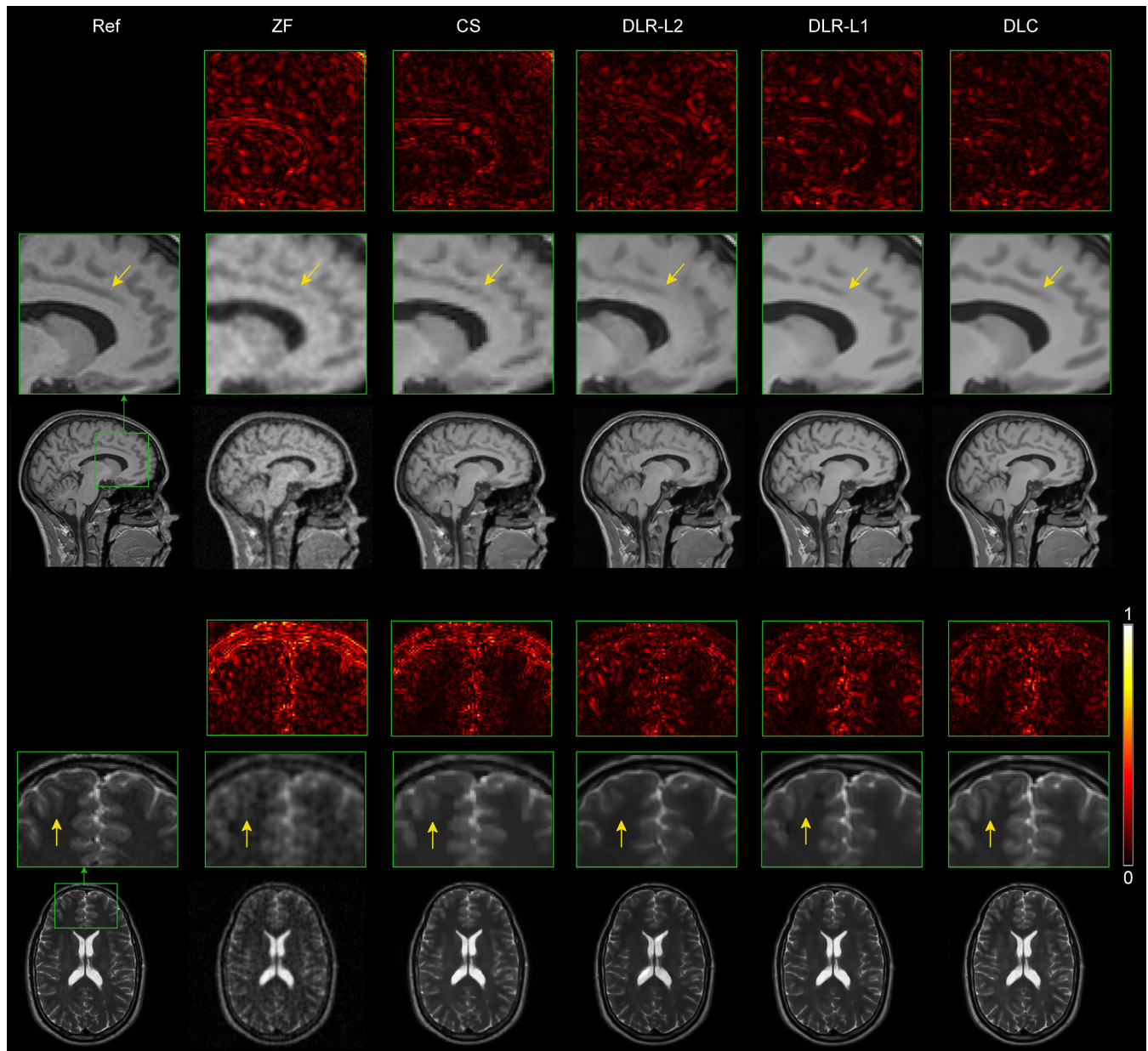


FIGURE 6. Results of the image reconstructions for T1 and T2 weighted images using the zero filled (ZF), the compressed sensing (CS), the DL-regression with L2 loss (DLR-L2), the DL-regression with L1 loss (DLR-L1) and the DL-classification (DLC) reconstructions (computed using equations (10, 11)). The enlarged images from each panel demonstrate that the image contrast and edges are better preserved in the DL classification network images compared to the CS and DL-regression reconstructed images. The small structure pointed out by yellow arrow shows that the delineation of the white matter and gray matter is more prominent in the DLC reconstructed images. The edge preservation can also be appreciated from the error images (scaled by 4x).

patch like artifacts in the reconstructed images (Figure 7, CS column). The DL regression network reconstructions performed better than the CS reconstructions but suffered from substantial loss of low contrast features in the presence of noise in the reconstructed images (Figure 7, DLR-L2 and DLR-L1). The DLC classification network performed substantially better than the other reconstruction methods, was able to reconstruct images with sharper edges and effectively removed the effect of noise in the reconstructed images. The results shown in Tables 3 and 4 demonstrate the quantitative

scores for the reconstructions in the presence of noise for T1 and T2 weighted images, respectively.

The DLC network was least affected by the injection of noise. For T1 weighted images, the SSIM decreased by 3.3 % for DLC while for DLR-L1, DLR-L2 and CS it decreased by 4 %, 4.1 % and 5.1 %, respectively.

Figure 8 shows the images reconstructed at an acceleration factor of 3 for 1D under sampling on the T1 weighted images along with the associated SSIM indexes. The reconstructed images (zoomed views) demonstrate that the image contrast

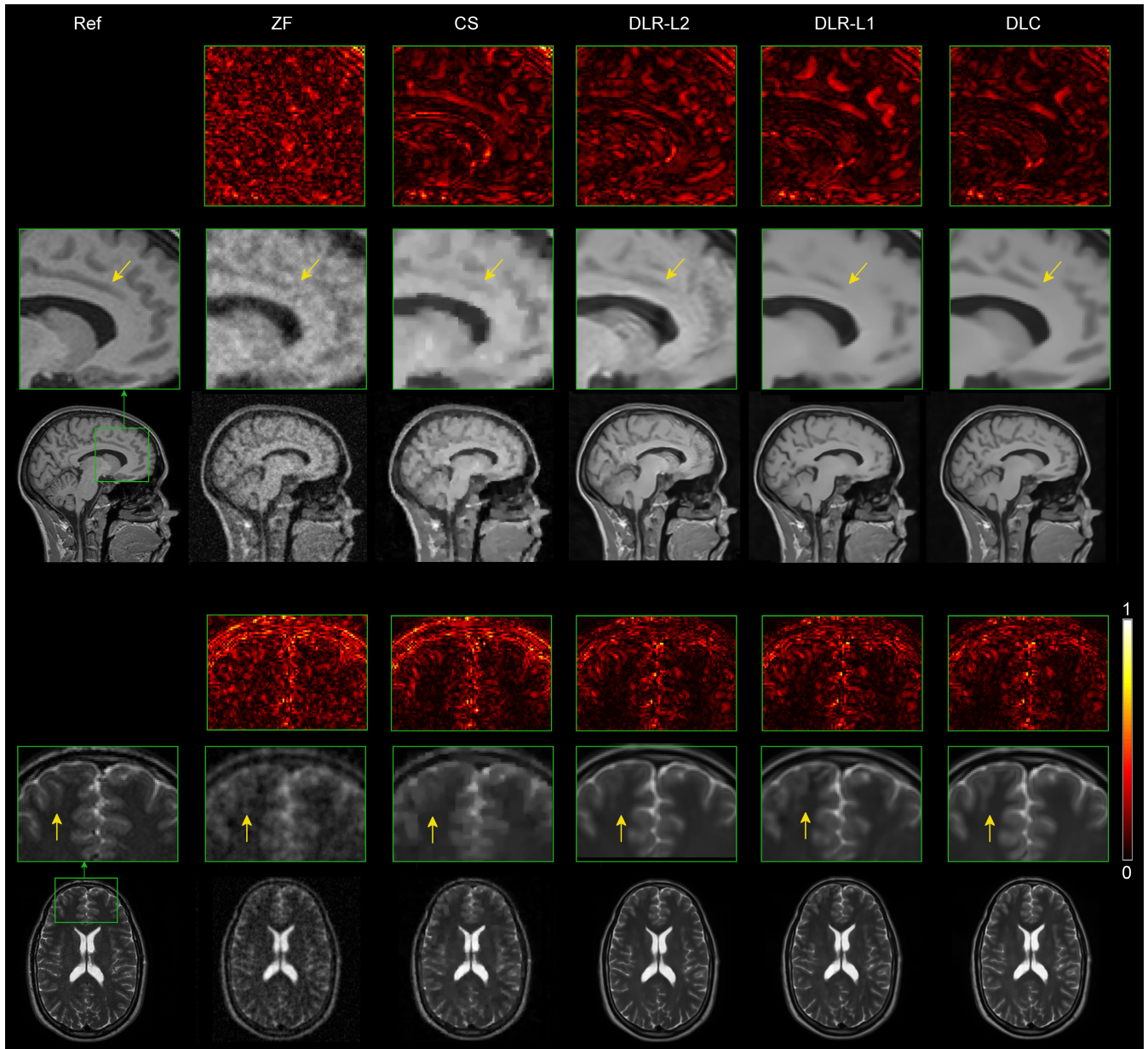


FIGURE 7. Results of the image reconstructions for T1 and T2 weighted images in the presence of added Gaussian noise using the zero filled (ZF), the Compressed Sensing (CS), the DL-regression with L2 loss (DLR-L2), the DL-regression with L1 loss (DLR-L1) and the DL-classification (DLC) reconstructions (computed using equations (10, 11)). The enlarged images from each panel demonstrate that the image contrast and edges are better preserved in the DL classification network images compared to the CS and DL-regression reconstructed images. The small structure pointed out by yellow arrow shows that the delineation of the white matter and gray matter is more prominent in the DLC reconstructed images. The edge preservation can also be appreciated from the error images (scaled by 4x).

and edges are better preserved in the DL classification network images compared to the CS and DL-regression images. The yellow arrow indicates the region of loss of contrast in the DLR-L2 images.

Additionally, we compared the performance of DLC network with ADMM-Net [34] for radial trajectory (Figure 9 (e)). Both the DLC and ADMM-Net were trained for acceleration factor of 5 (sampling ratio = 20%). The ADMM-Net was trained in MATLAB using the code from (<https://github.com/yangyan92/Deep-ADMM-Net>). Figure 9 shows images reconstructed for the acceleration factor of 5

TABLE 3. Quantitative scores for T1 weighted images in the presence of noise.

Methods	SSIM	PSNR	Relative Error	MSE
Compressed Sensing	0.8909	28.02	0.0753	1.577e-3
DLR-L2	0.9077	28.76	0.0692	1.330e-3
DLR-L1	0.9120	28.83	0.0686	1.308e-3
DL Classification	0.9227	30.14	0.0641	1.114e-3

along with the SSIM index which was 0.9309 for ADMM-Net and was 0.9529 for DLC. The image reconstructed using DLC (Figure 9 (d)) was visually more pleasing and the

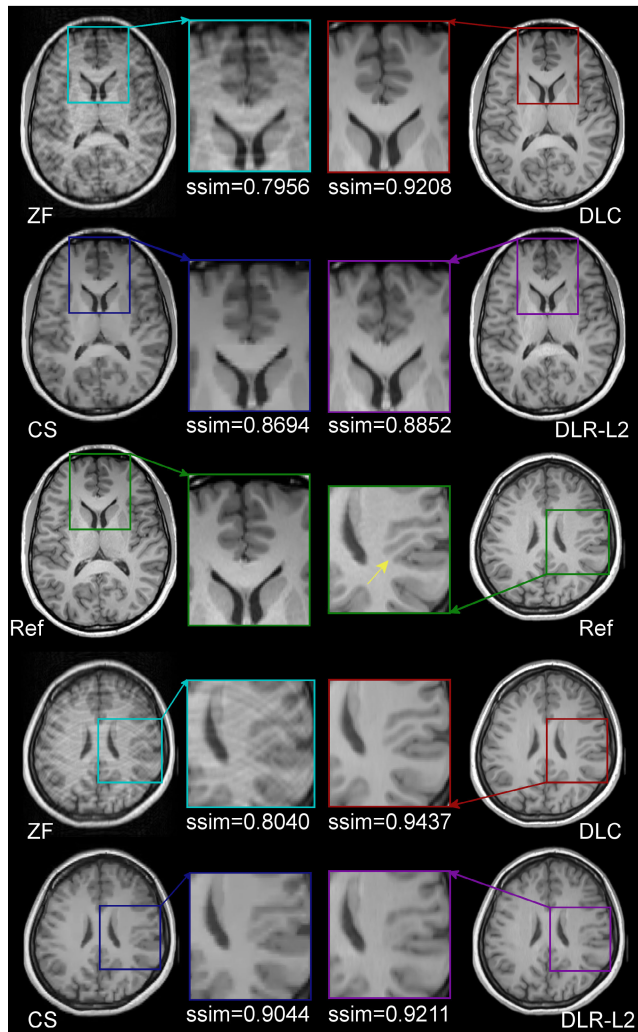


FIGURE 8. Results of the image reconstructions for two axial sections of T1 weighted MPAGE for the acceleration factor of 3 using the zero filled (ZF), the Compressed Sensing (CS), the DL-regression (DLR-L2) and the DL-classification (DLC) reconstructions along with SSIM indexes. The enlarged images from each panel demonstrate that the image contrast and edges are better preserved in the DL classification network images compared to the CS and DL-regression images.

TABLE 4. Quantitative scores for T2 weighted images in the presence of noise.

Methods	SSIM	PSNR	Relative Error	MSE
Compressed Sensing	0.9306	29.11	0.0679	1.224e-3
DLR-L2	0.9590	31.66	0.0506	6.813e-4
DLR-L1	0.9561	30.84	0.0556	8.231e-4
DL Classification	0.9615	31.78	0.0499	6.338e-4

small structures are visible clearly. In contrast, the image reconstructed from the ADDM-Net (Figure 9 (c)) consisted of patchy artefacts and small structures were not visible clearly.

The trained networks (noise-free trained) on the T1 weighted images were further tested on the brain tumor images with a BRATS segmentation dataset [31]. Both the DLC and DLR-L1 networks were directly used on the brain

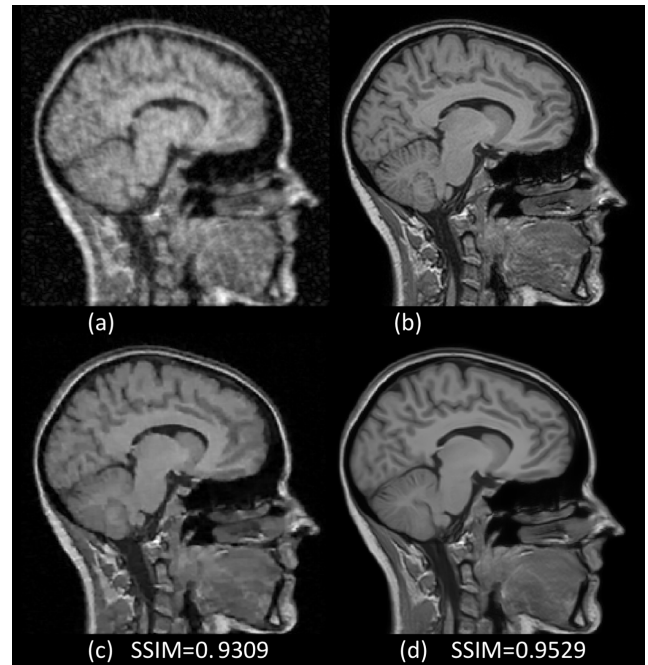


FIGURE 9. Results of image reconstruction for an image with radial sampling pattern for acceleration factor of five with ADMM-Net and classification network DLC. (a): Zero filled image for acceleration factor of five; (b) fully sampled reference image; (c) image reconstructed from ADMM-Net; (d) image reconstructed from DLC. The bottom of the panel (c-d) shows the SSIM index.

TABLE 5. Quantitative scores for brain tumor image.

Methods	SSIM	PSNR	Relative Error	MSE
DLR-L1	0.9802	36.93	0.0270	2.026e-4
DL Classification	0.9822	37.30	0.0258	1.862e-4

tumor images without any further training or fine-tuning. These networks have never seen the images with tumors. For the acceleration factor of eight, it was observed that the images (Figure 10) reconstructed with the DLC network provide better contrast compared to the images reconstructed with DLR-L1, which suffered from smoothening and contrast wash. The edges of the tumor were better preserved in the DLC reconstructions compared to the DLR-L1 reconstruction, as seen in the error images in Figure 10 and the quantitative results shown in Table 5.

We computed the inference time for both the regression and the classification networks using a Nvidia Tesla V100 GPU. The inference time for a single slice of 256×256 was approximately 0.36 seconds for both the networks. We further analyzed the probability distribution predicted by the DLC network across different regions of the image. The predicted probability distributions for two different pixels in a DL classification reconstructed image are shown in Figure 10. For the pixels falling in the region of relatively flat texture, the confidence was high (Figure 11 (b)) and the probability distribution was narrow. For the pixels in regions of high

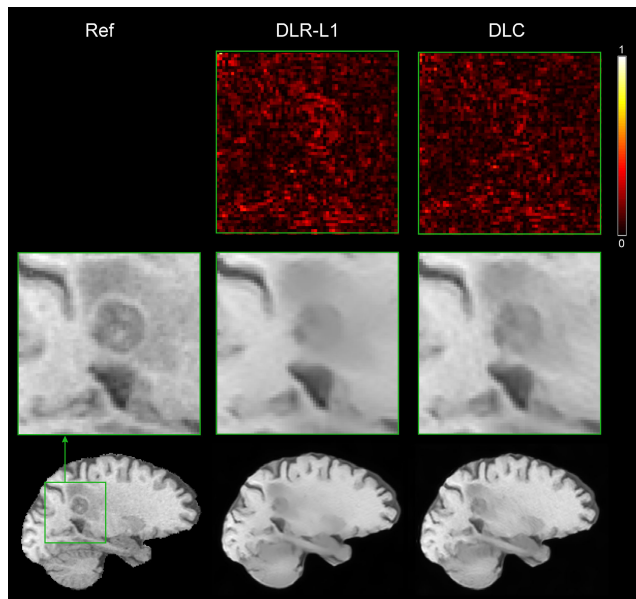


FIGURE 10. Results of image reconstruction for an image with brain tumor for acceleration factor of eight with regression network trained with L1 loss (DLR-L1) and classification network trained with cross entropy loss (DLC). Ref: Fully sampled reference image; DLR-L1: image reconstructed with regression network trained on L1 loss; DLC: image reconstructed with classification network. The edges of the tumor were preserved better in the DLC reconstructions compared to DLR-L1 reconstruction as seen in the error images.

variance in the image, which is near the sharp edges, the confidence was comparatively low and the probability distribution was wide (Figure 11 (c)). As the probability distribution was symmetrical, the widening of the probability distribution does not result in a change in the mean value, and thus, the network was able to recover more accurate intensities of the pixels near the sharp edges. In contrast, the regression network was not able to effectively differentiate between the adjacent pixels at the edges, which resulted in blurring of the edges.

V. SUMMARY AND DISCUSSION

In order to further improve the existing deep learning image reconstruction methods, in this work, we have developed a novel framework for transforming image reconstruction into pixel classification. Conventionally, image quantization occurs during digital image archiving (e.g. 16-bit DICOM image storage). In this work, we introduce a framework that directly reconstructs “digital” images. The proposed method can be used with a multitude of deep learning reconstruction methods. The method was extensively tested on T1 and T2 weighted images and compared with conventional compressed sensing and DL regression network image reconstructions. The images reconstructed from the proposed method outperformed the images reconstructed from the other methods, both qualitatively and quantitatively.

We have applied and validated the proposed method using T1 and T2 weighted MR images, and with T1 tumour contrast images. These anatomical images represent a high contrast

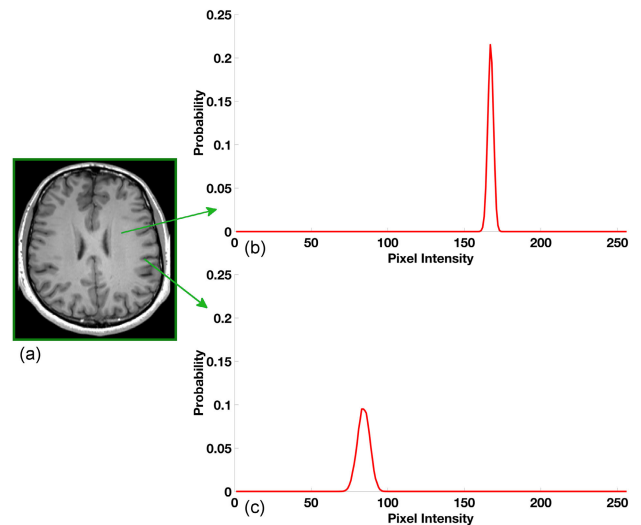


FIGURE 11. (a): DL classification network reconstructed image; (b-c): the predicted probability distribution by DL classification network at pixels pointed out by the white arrow.

and a highly dynamic range in the family of MR images. With the successful validation using these datasets, we can reasonably expect that the proposed method will work well with other contrasts and dynamic range, such as functional MRI, perfusion and diffusion MRI.

The deep learning methods use the process of training to find an underlying model relating the image with undersampling artifacts and the image without undersampling artifacts. The DL regression networks find an exact model relating the artifact image to the clean image. However, in the proposed DL classification network, we successfully applied the concept of image segmentation to solve the image reconstruction regression problem. The application of a classification approach to a regression problem turns out to be better at preserving the low contrast features of the images, as evident from the results in Figures 6-10. Notably, the classification model performs very well when noise is present in the image and was able to remove the artifacts from the noisy images with undersampling, as evident from the results shown in Figures 6 and 7.

The proposed approach is generic in nature and demonstrates that many of the advances made in deep learning classification problems can be integrated into MR image reconstruction tasks. One example is the problem of class imbalance during the training of a deep learning model. Class imbalance is a well-studied problem in deep learning for object detection. Class imbalance issue is also present in MR images, as most of the background pixels are zero. Therefore, class imbalance coping methods, such as focal loss [33], can also be integrated with the proposed classification approach.

It is reasonable to expect the quantization step to perform well with a variety of image contrasts, including proton density. More advanced quantization algorithms, such as non-uniform quantizer, can also be used in conjunction with the DL-C network. Furthermore, for complex image

reconstruction, including susceptibility-weighted imaging (SWI) and phase-contrast imaging, it is possible to use two channels in the DL network - one for the real component and another for the imaginary component or magnitude and phase. The phase values range from 0 to 2π , and using an 8-bit DLC network the precision will be $\frac{2\pi}{256}$, while for a 16-bit network, the precision will be $\frac{2\pi}{65536}$, resulting in very high precision. Another issue with phase image reconstruction is the background phase signal caused by B0 field inhomogeneity. The background phase varies from scan to scan and subject-to-subject; therefore, a background phase correction method is likely needed in combination with the DL framework.

One limitation of using a pixel classification network is an exponential increase in the memory requirement as the number of quantization bits are increased. If the required large memory is available, the proposed method can be easily extended to higher bit depth by increasing the number of output channels in the DLC network. However, this may result in overfitting due to the increased number of network parameters. In order to circumvent this limitation, we proposed a novel divide-and-conquer approach (section I-E) for extending the classification network to predict higher bit depth pixel values without any significant increase in the number of network parameters. For a 16-bit classification network, the number of parameters was only increased by 0.6% compared to an 8-bit network.

For prospective undersampling where the data is inherently multi-channel, the undersampled data can be combined to a single image which can be processed through the DL network. However, training a separate network on a multi-channel image as an input and a coil combined image as an output would result in better performance.

We performed experiments on the effect of the sampling pattern at the time of inference. The performance of the network trained on an acceleration factor of 8 was degraded when the inference was performed on the acceleration factor of 10. However, the effect on performance was marginal for a fixed acceleration factor change in the sampling pattern drawn from the same Gaussian distribution.

Furthermore, using a non-uniform quantizer can further improve the encoding efficiency of information during quantization step and the performance of the DL network. The framework introduced in this paper has the potential to leverage great advances in this area of research.

VI. CONCLUSION

A generic framework for transforming image reconstruction into pixel classification, which can be used with many deep learning based image reconstruction methods is demonstrated. The proposed method restores low contrast features better than the other standard methods. The method is robust to noise and can reconstruct high-contrast images in relatively low signal to noise ratio scenarios.

REFERENCES

- [1] A. KrizhKrizhevsky, I. Sutskever, and G. E. Hinton, "ImageNet classification with deep convolutional neural networks," in *Proc. Adv. Neural Inf. Process. Syst.*, I. S. Alex and G. E. Hinton, 2012, pp. 1–9.
- [2] P. Mansfield, "Multi-planar image formation using NMR spin echoes," *J. Phys. C, Solid State Phys.*, vol. 10, no. 3, p. 55, 1977.
- [3] J. Frahm, A. Haase, and D. Matthaei, "Rapid NMR imaging of dynamic processes using the FLASH technique," *Magn. Reson. Med.*, vol. 3, no. 2, pp. 321–327, 1986.
- [4] Z. Chen, J. Zhang, R. Yang, P. Kellman, L. A. Johnston, and G. F. Egan, "IIR GRAPPA for parallel MR image reconstruction," *Magn. Reson. Med.*, vol. 63, no. 2, pp. 502–509, 2010.
- [5] M. A. Griswold, P. M. Jakob, R. M. Heidemann, M. Nittka, V. Jellus, J. Wang, B. Kiefer, and A. Haase, "Generalized autocalibrating partially parallel acquisitions (GRAPPA)," *Magn. Reson. Med.*, vol. 47, no. 6, pp. 1202–1210, 2002.
- [6] K. P. Pruessmann, M. Weiger, M. B. Scheidegger, and P. Boesiger, "SENSE: Sensitivity encoding for fast MRI," *Magn. Reson. Med.*, vol. 42, no. 5, pp. 952–962, 1999.
- [7] D. K. Sodickson, "Tailored SMASH image reconstructions for robust *in vivo* parallel MR imaging," *Magn. Reson. Med.*, vol. 44, no. 2, pp. 243–251, 2000.
- [8] D. K. Sodickson and W. J. Manning, "Simultaneous acquisition of spatial harmonics (SMASH): Fast imaging with radiofrequency coil arrays," *Magn. Reson. Med.*, vol. 38, pp. 591–603, 1997.
- [9] M. Lustig, D. Donoho, and J. M. Pauly, "Sparse MRI: The application of compressed sensing for rapid MR imaging," *Magn. Reson. Med.*, vol. 58, no. 6, pp. 1182–1195, 2007.
- [10] M. Lustig, D. L. Donoho, J. M. Santos, and J. M. Pauly, "Compressed sensing MRI," *IEEE Signal Process. Mag.*, vol. 25, no. 2, pp. 72–82, Mar. 2008.
- [11] C. A. Baron, N. Dwork, J. M. Pauly, and D. G. Nishimura, "Rapid compressed sensing reconstruction of 3D non-Cartesian MRI," *Magn. Reson. Med.*, vol. 79, no. 5, pp. 2685–2692, 2018.
- [12] M. Murphy and M. Lustig, " ℓ_1 minimization in ℓ_1 -SPIRiT compressed sensing MRI reconstruction," in *Proc. GPU Comput. Gems Emerald Ed.*, 2011, pp. 723–735.
- [13] T. M. Quan, S. Han, H. Cho, and W. K. Jeong, "Multi-GPU reconstruction of dynamic compressed sensing MRI," in *Medical Image Computing and Computer-Assisted Intervention—MICCAI* (Lecture Notes in Computer Science), vol. 9351, N. Navab, J. Hornegger, W. Wells, and A. Frangi, Eds. Cham, Switzerland: Springer, 2015.
- [14] M. Zimmermann, Z. Abbas, K. Dzieciol, and N. J. Shah, "Accelerated parameter mapping of multiple-echo gradient-echo data using model-based iterative reconstruction," *IEEE Trans. Med. Imag.*, vol. 37, no. 2, pp. 626–637, Feb. 2018.
- [15] Y. A. LeCun, Y. Bengio, and G. E. Hinton, "Deep learning," *Nature*, vol. 521, no. 7553, pp. 436–444, 2015.
- [16] G. Litjens, T. Kooi, B. E. Bejnordi, A. A. A. Setio, F. Ciompi, M. Ghafoorian, J. A. W. M. van der Laak, B. van Ginneken, and C. I. Sánchez, "A survey on deep learning in medical image analysis," *Med. Image Anal.*, vol. 42, pp. 60–88, Dec. 2017.
- [17] K. Pawar, Z. Chen, N. J. Shah, and G. Egan, "Residual encoder and convolutional decoder neural network for glioma segmentation," in *Proc. Int. MICCAI Brainlesion Workshop*, 2018, pp. 263–273.
- [18] O. Ronneberger, P. Fischer, and T. Brox, "U-Net: Convolutional networks for biomedical image segmentation," in *Proc. MICCAI*, 2015, pp. 234–241.
- [19] E. Shelhamer, J. Long, and T. Darrell, "Fully convolutional networks for semantic segmentation," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 39, no. 4, pp. 640–651, Apr. 2017.
- [20] K. Hammernik, T. Klatzer, E. Kobler, M. P. Recht, D. K. Sodickson, T. Pock, and F. Knoll, "Learning a variational network for reconstruction of accelerated MRI data," *Magn. Reson. Med.*, vol. 79, no. 6, pp. 3055–3071, 2018.
- [21] K. H. Jin, M. T. McCann, E. Froustey, and M. Unser, "Deep convolutional neural network for inverse problems in imaging," *IEEE Trans. Image Process.*, vol. 26, no. 9, pp. 4509–4522, Sep. 2017.
- [22] D. Lee, J. Yoo, and J. C. Ye, "Deep residual learning for compressed sensing MRI," in *Proc. IEEE 14th Int. Symp. Biomed. Imag.*, Apr. 2017, pp. 15–18.
- [23] T. M. Quan, T. Nguyen-Duc, and W.-K. Jeong, "Compressed sensing MRI reconstruction using a generative adversarial network with a cyclic loss," *IEEE Trans. Med. Imag.*, vol. 37, no. 6, pp. 1488–1497, Jun. 2018.

- [24] J. Schlemper, J. Caballero, J. V. Hajnal, A. Price, and D. Rueckert, "A deep cascade of convolutional neural networks for MR image reconstruction," in *Proc. Inf. Process. Med. Imag.*, 2017, pp. 647–658.
- [25] S. Wang, Z. Su, L. Ying, X. Peng, S. Zhu, F. Liang, D. Feng, and D. Liang, "Accelerating magnetic resonance imaging via deep learning," in *Proc. IEEE 13th Int. Symp. Biomed. Imag. (ISBI)*, Apr. 2016, pp. 514–517.
- [26] G. Yang, S. Yu, H. Dong, G. Slabaugh, P. L. Dragotti, X. Ye, F. Liu, S. Arridge, J. Keegan, Y. Guo, D. Firmin, "DAGAN: Deep de-aliasing generative adversarial networks for fast compressed sensing MRI reconstruction," *IEEE Trans. Med. Imag.*, vol. 37, no. 6, pp. 1310–1321, Jun. 2018.
- [27] B. Zhu, J. Z. Liu, S. F. Cauley, B. R. Rosen, and M. S. Rosen, "Image reconstruction by domain-transform manifold learning," *Nature*, vol. 555, no. 7697, pp. 487–492, 2018.
- [28] J. C. Ye, Y. Han, and E. Cha, "Deep convolutional framelets: A general deep learning framework for inverse problems," *SIAM J. Imag. Sci.*, vol. 11, no. 2, pp. 991–1048, 2018.
- [29] J. F. Blinn, "Quantization error and dithering," *IEEE Comput. Graph. Appl.*, vol. 14, no. 4, pp. 78–82, Jul. 1994.
- [30] R. M. Gray and D. L. Neuhoff, "Quantization," *IEEE Trans. Inf. Theory*, vol. 44, no. 6, pp. 2325–2384, Oct. 1998.
- [31] B. H. Menze et al., "The multimodal brain tumor image segmentation benchmark (BRATS)," *IEEE Trans. Med. Imag.*, vol. 34, no. 10, pp. 1993–2024, Oct. 2015.
- [32] B. Widrow, "Statistical analysis of amplitude-quantized sampled-data systems," *Amer. Inst. Electr. Engineers, II, Appl. Ind., Trans.*, vol. 79, no. 6, pp. 555–568, 1961.
- [33] T.-Y. Lin, P. Goyal, R. Girshick, K. He, and P. Dollar, "Focal loss for dense object detection," in *Proc. IEEE Int. Conf. Comput. Vis.*, Jun. 2017, pp. 2980–2988.
- [34] J. Sun, H. Li, and Z. Xu, "Deep ADMM-Net for compressive sensing MRI," in *Proc. Adv. Neural Inf. Process. Syst.*, 2016, pp. 10–18.



algorithms on SIMD embedded processor and development of instruction set architecture for SIMD image processor. He is currently working with Monash Biomedical Imaging, Melbourne, Australia. His passion includes the development of novel imaging technologies. His expertise includes MR physics, MR pulse sequence design and MR image reconstruction, deep learning, and computer vision.

KAMLESH PAWAR received the Ph.D. degree in electrical and computer system engineering on the development of novel imaging methodologies and reconstruction methods for accelerated MR imaging, jointly conferred by IIT Bombay, India, and Monash University, Australia, in 2014. He was offered Tata Consultancy Services (TCS) Research Scholar Fellowship for his Ph.D. degree. From 2014 to 2016, he worked at Cadence Design Systems on the implementation of computer vision



Monash Biomedical Imaging, in 2014, and since then, he has been leading the Imaging Analysis team. He has authored more than 60 peer-reviewed journal articles and conference papers. He has invented ten patents and many of them have been commercialized and applied in clinical practice. His current research interests include MRI and PET data acquisition, image reconstruction and artifact reduction using machine learning, and image analysis for neuroscience and clinical applications. He was a recipient of the Douglas Lampard Research Prize and received Medal from Monash University. He serves as a Programme Committee Member for the International Society of Magnetic Resonance in Medicine (ISMRM) and an Academic Editor for PLoS One.

ZHAOLIN CHEN received the B.Eng. degree in automation from the Harbin Institute of Technology, in 2002, and the Ph.D. degree in MRI image reconstruction from the Department of Electrical and Computer Systems Engineering, Monash University, in 2007. From 2008 to 2010, he was a Postdoctoral Fellow with the Florey Neuroscience Institutes and the University of Melbourne. From 2010 to 2014, he was an Imaging Scientist with Philips Healthcare, The Netherlands. He joined



spanning from the basic conception and design of new MRI and MR-PET methodology through to its implementation and application. He is also the Director of the Institute of Neuroscience and Medicine-4 (Medical Imaging Physics) and the Institute of Neuroscience and Medicine-11 (Molecular Neuroscience and Neuroimaging), Research Centre, Jülich. He is also a Professor of MRI physics with the Department of Neurology, University Hospital, Aachen. He was awarded the Toshiba Fellowship which took him to Japan, where he worked on the development of methods for magnetic resonance imaging and spectroscopy. Thereafter, he was awarded the Herschel Smith Fellowship with the University of Cambridge. He awarded a prestigious three year VESKI Fellowship.

N. JON SHAH received the Ph.D. degree from the University of Manchester. From 2015 to 2017, he was a Distinguished Professor with the Monash Institute of Medical Engineering (MIME). During his three year stay, he carried out research in basic MR methodology and translated those methods into Addenbrooke's Hospital. He moved to Germany to work in the area of cardiac MRI. He is currently with the Research Centre, Jülich, where he is responsible for a Research Programme



ology, neuroimaging, informatics, neural modeling, and neuro engineering. For the past two decades, he has been one of the Australia's pioneering and internationally recognized neuroimaging and neuroscience researchers, using PET and MRI in human and animal model neuroscience research, and establishing the emerging field of neuroinformatics in Australia.

GARY. F. EGAN is currently a Distinguished Professor and the Foundation Director of the Monash Biomedical Imaging (MBI) research facilities, Monash University, and the Director of the Australian Research Council Centre of Excellence for Integrative Brain Function. The Centre's research focus is to understand the link between brain activity and human behavior by integrating research across Australia's leading brain researchers in the fields of anatomy, physiology, neuroimaging, informatics, neural modeling, and neuro engineering.

...