



Correcting synthetic MRI contrast-weighted images using deep learning

Sidharth Kumar^{a,*}, Hamidreza Saber^{b,c}, Odelin Charron^b, Leorah Freeman^{b,d}, Jonathan I. Tamir^{a,d,e}

^a Chandra Family Department of Electrical and Computer Engineering, The University of Texas at Austin, Austin 78712, TX, USA

^b Dell Medical School Department of Neurology, The University of Texas at Austin, Austin 78712, TX, USA

^c Dell Medical School Department of Neurosurgery, The University of Texas at Austin, Austin 78712, TX, USA

^d Dell Medical School Department of Diagnostic Medicine, The University of Texas at Austin, Austin 78712, TX, USA

^e Oden Institute for Computational Engineering and Sciences, The University of Texas at Austin, Austin 78712, TX, USA

ARTICLE INFO

Keywords:

Synthetic MRI
Physics-enabled deep learning
Multi-contrast MRI
Quantitative MRI

ABSTRACT

Synthetic magnetic resonance imaging (MRI) offers a scanning paradigm where a fast multi-contrast sequence can be used to estimate underlying quantitative tissue parameter maps, which are then used to synthesize any desirable clinical contrast by retrospectively changing scan parameters in silico. Two benefits of this approach are the reduced exam time and the ability to generate arbitrary contrasts offline. However, synthetically generated contrasts are known to deviate from the contrast of experimental scans. The reason for contrast mismatch is the necessary exclusion of some unmodeled physical effects such as partial voluming, diffusion, flow, susceptibility, magnetization transfer, and more. The inclusion of these effects in signal encoding would improve the synthetic images, but would make the quantitative imaging protocol impractical due to long scan times. Therefore, in this work, we propose a novel deep learning approach that generates a multiplicative correction term to capture unmodeled effects and correct the synthetic contrast images to better match experimental contrasts for arbitrary scan parameters. The physics inspired deep learning model implicitly accounts for some unmodeled physical effects occurring during the scan. As a proof of principle, we validate our approach on synthesizing arbitrary inversion recovery fast spin-echo scans using a commercially available 2D multi-contrast sequence. We observe that the proposed correction visually and numerically reduces the mismatch with experimentally collected contrasts compared to conventional synthetic MRI. Finally, we show results of a preliminary reader study and find that the proposed method statistically significantly improves in contrast and SNR as compared to synthetic MR images.

1. Introduction

Magnetic resonance imaging (MRI) is an important non-invasive clinical imaging modality that does not use ionizing radiation. A big benefit of MRI is its ability to capture a multitude of tissue contrasts by changing the acquisition parameters, providing complementary information to characterize and assess pathology. A typical clinical MRI protocol consists of multiple independent scans, contributing to an overall lengthy exam [1]. Long exams lead to high associated costs for patients and also reduce the overall throughput of the scanner. Moreover, it is difficult for patients (especially pediatric and elderly) to hold still for a long scanning period, hence making the images susceptible to motion-induced artifacts [2]. Even when the scans are artifact-free, the contrast weightings represent a qualitative signal intensity that must be

visually inspected and reported on by a clinical radiologist, and thus only qualitative signal abnormalities can be reported.

Recently there has been substantial interest in extracting quantitative information from MRI in addition to conventional qualitative images [3,4]. In quantitative MRI, biophysical tissue and system parameters are estimated using specialized multi-contrast acquisitions that acquire the MRI signal at multiple contrast points. While conventional quantitative imaging is slow in nature, recent work has led to fast and comprehensive multi-contrast scans [3–6]. After tissue quantification, arbitrary contrasts can be synthetically created by evaluating the MRI signal equation for a retrospectively chosen set of scan parameters in silico. A large body of work has focused on accurate quantification of T1, T2, and proton density (PD) maps [4,7]. System imperfections including magnetic and radio-frequency (RF) field inhomogeneities

* Corresponding author.

E-mail address: sidharth.kumar@utexas.edu (S. Kumar).

<https://doi.org/10.1016/j.mri.2023.11.015>

Received 30 August 2023; Received in revised form 30 November 2023; Accepted 30 November 2023

Available online 12 December 2023

0730-725X/© 2023 Elsevier Inc. All rights reserved.

have also been incorporated [3,8]. However, the underlying MRI physics are more nuanced and have many confounding parameters that contribute to the final image contrast, including magnetization transfer, partial voluming, diffusion, and susceptibility, to name a few. While it is possible to extend the quantitative imaging protocol to map these additional effects, the scan time can become impractical, and the physical modeling can be challenging. As a result, the synthetic contrast image often has subtle differences when compared to the image obtained from a real acquisition. For example, synthetic T2-FLAIR contrast is known to suffer from hyperintense signal artifacts surrounding the cerebrospinal fluid (CSF) [9,10], likely due to the over-simplification of the MRI signal model.

To combat contrast-mismatch effects, researchers have recently leveraged deep learning to correct these artifacts [11–13]. These methods pose contrast correction as a supervised image-to-image translation, in which a particular set of experimental contrasts are acquired and treated as reference images, and a deep neural network is trained to map the incorrect contrast to the experimental contrast. While the results are striking, the approaches are limited in that they only correct a particular set of contrasts that are collected at training time, and therefore they lose the ability to synthesize arbitrary contrast-weighted images offline.

Therefore in this work, we propose a novel physics-enabled deep learning method to correct arbitrary contrast-weighted images generated by synthetic MRI. While we also pose the problem as supervised image translation, we instead aim to learn the mapping from synthetic MRI to experimental contrast as a function of scan parameters. Motivated by the main unmodeled effects in fast spin-echo (FSE) imaging, our proposed method generates a spatially dependent multiplicative correction term and a scaled apparent longitudinal relaxation time (T_1). We provide several different experimental contrast-weighted images as training samples so that our network implicitly learns the relationship between scan parameters and the physical effects that are not captured by the simplified signal equation. We accomplish this using a conditional generative adversarial network (cGAN) framework complimented with perceptual loss from a separate pre-trained network [14–16]. To incorporate scan parameters, we include quantitative maps and synthetic contrast images along with scan parameters as additional channels for the generator network. While we demonstrate our approach for a specific multi-contrast sequence available on our scanner, we emphasize that the framework could be used for other sequences such as magnetic resonance fingerprinting (MRF) [3].

The main contributions of our paper are as follows: i) we propose a novel deep learning method that explicitly incorporates the scan parameters to correct arbitrary synthetic MRI contrast-weighted images derived from a multi-contrast sequence; ii) we show how such models can be trained using a 2D multi-contrast sequence as proof of principle, so that the pipeline can be used for other custom multi-contrast sequences; iii) we evaluate our method on subjects and MRI contrasts that were not included in the training set to give evidence that the approach implicitly accounts for unmodeled physical effects. The proposed model gives better numerical error metrics than the direct and residual correction model. We will make our data and implementation available publicly upon publication.

1.1. Related works

Image-to-image translation has been highly successful in computer vision, in which conditional generative models are trained to map from one image style to an output image style [14,17]. Analogous to these tasks, researchers in medical imaging have also proposed image-to-image translation models; for example, computed tomography (CT) to MRI [18], positron emission tomography (PET) to CT [19], medical image segmentation [20], low-to-high gadolinium dose contrast-enhanced brain MRI [15], etc. Specifically in MRI, translating from one MRI contrast-weighted image to another contrast-weighted image is

a well investigated problem. Authors in [16] propose multi-contrast synthesis through cGANs and demonstrate the applicability by translating from T1-weighted to T2-weighted images and vice versa. A multi-stream approach was also proposed to join information from one-to-one and many-to-one translation streams using a fusion block [21]. The work in [22] proposes MRI motion correction through cGANs by translating from motion-corrupted to motion-free images by incorporating FSE acquisition dynamics. Authors in [23] proposed a multi-input, multi-output GAN network to generate missing MRI sequences using the redundant information from other available sequences. The work in [24] proposed a hybrid-fusion network to generate target MRI contrasts from source images. The overall network consisted of two small sub-networks, the first network to learn representations from each input modality and the second network to fuse the common latent representation and synthesize target images. The authors in [25], proposed an edge-aware network that captures the textural details of MR images to improve the overall final image quality in cross-modality MR synthesis. A novel method to generate 3D brain MRI from learned representations using variational auto-encoder and GAN was proposed in [26] that generates high-quality image data from limited training data. Authors in [27] applied the cGAN framework to reconstruct patient faces from anonymized T1 sagittal slices using unsupervised training and were able to recover patient information from both face blurred and face removed data. Similarly, work in [28] showed the multimodal image synthesis using GANs on glioma patients. Authors in [29] proposed DiamondGAN to do non-aligned cross-modality synthesis and performed a radiologist evaluation study to show that trained radiologists were not able to differentiate between experimental and synthetic MRI images.

The other MRI translation task frequently considered is to correct the synthetic MRI contrast. The authors in [13] looked at direct contrast synthesis from a multi-contrast scan using a temporal-convolutional network on a per-pixel basis. The work in [11] extended the direct contrast synthesis to operate on the whole input image and output the desired contrast using cGANs whereas work in [30] presented a convolutional encoder-decoder network to directly generate multiple contrast image from base multi-echo sequence. Authors in [31] proposed a deep learning method to improve the T2-FLAIR contrast generation directly from base multi-contrast images. On the same track, the work in [12] corrected the synthetic contrasts of T2-FLAIR on a per-pixel basis using convolutional neural networks. Similarly, authors in [10] used a generative network to translate directly from the echo images of the base multi-contrast images to the FLAIR contrasts. The review paper [32] discusses the synthetic MRI methods to generate multiple contrast images along with some of the limitations in synthetic contrast generation; in particular, they report lower quality in synthetic FLAIR images. Common to all these works is the need to collect the ground-truth experimental contrasts that are desired, and the restriction to only correcting those contrasts. In comparison, our goal is to maintain the ability to synthesize contrasts corresponding to arbitrary scan parameters. All of the discussed works follow the approximate framework shown in Fig. 1 to directly translate from base multi-contrast images to final target contrasts, which deviate from the premise of synthetic MRI to generate arbitrary contrasts.

Several clinical validation studies have shown the benefits of synthetic MRI, while also highlighting its limitations. In [33], clinicians rated all synthetic contrasts to be inferior to the conventional scans except T_2 weighted images. Prior to that, authors in [34] conducted a clinical validation study using the Multi Delay Multi Echo (MDME) sequence. It was reported that overall image quality was similar for all contrasts except for FLAIR where a conventional scan was still clinically necessary. Similarly, work in [35] reported satisfactory image quality for all contrasts except T2 FLAIR. The thread of these works is that the inversion recovery contrast images were hard to correctly synthesize, with implications on their clinical use.

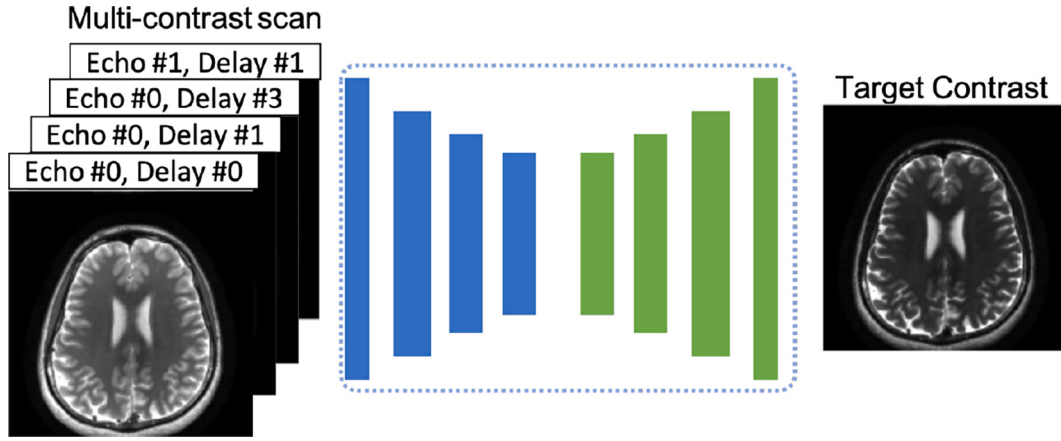


Fig. 1. Conventional contrast correction using deep neural networks, the objective is to correct a particular contrast of interest or to translate from the base multi-contrast images to a fixed contrast of interest.

2. Methods

2.1. Synthetic MRI

The premise of synthetic MRI is to use a multi-contrast sequence to estimate underlying tissue and system parameter maps and then generate new contrasts in silico. For this work, we specialize our exposition to the 2D MDME sequence [4], which is capable of mapping proton density (PD), longitudinal (T1), and transverse (T2) relaxation. This choice is somewhat arbitrary and only used to illustrate our method in the experimental results section, but the core ideas presented are expected to work with other base multi-contrast sequences. The MDME sequence uses an FSE train at different echo time (TE) points to encode T₂ information. Additionally, different delay times (T_D) are set in between a saturation pulse and an acquisition block, which encodes the T₁ information of a particular slice. The underlying tissue parameters can then be estimated using the known signal model, for example through dictionary matching [3]. Finally, synthetic contrasts are generated by simulating the MRI signal equation (either analytically or algorithmically) for arbitrary scan parameters. The MDME signal after the saturation pulse is known to follow the following expression:

$$I = PD \left[\frac{1 - (1 - \cos\theta)e^{-\frac{TR}{T_1}} - \cos(\theta)e^{-\frac{TR}{T_1}}}{1 - \cos(\theta)\cos(\gamma)e^{-\frac{TR}{T_1}}} \right] e^{-\frac{TE}{T_2}}, \quad (1)$$

where PD is the proton density, γ is the RF excitation flip angle, θ is the RF saturation flip angle, T_D is the delay time, TE is the echo time, and TR is the repetition time. The specific protocol settings for the MDME sequence are provided in the later part of this section. To estimate parameter maps with dictionary fitting, a signal dictionary was created by varying the T₁ values from 100 to 6000 ms with a step size of 20 ms and the T₂ values from 10 to 1000 ms with a step size of 2 ms to cover relaxation parameters in the brain at 3 Tesla. The signal evolution for the dictionary generation was simulated using the Extended Phase Graph algorithm [36]. Dictionary matching was performed on the experimental signal to evaluate T1, T2, and PD values.

As a proof of concept, in this work, we restrict our method to synthesizing arbitrary inversion recovery spin-echo (IR-SE) contrasts as typically these synthetic contrast images are the most susceptible to artifacts [9,10,33–35]. The IR-SE contrast sequence can be described by $\{\beta_1 - TI - \beta_2 - \frac{TE}{2} - \beta_3 - (TR - TI - \frac{TE}{2})\}$, where typically $\beta_1 = \beta_3 = 180^\circ$ and $\beta_2 = 90^\circ$. Assuming no RF inhomogeneity, the IR-SE contrast can be generated according to [37].

$$I = PD \left[1 - 2e^{-\frac{TI}{T_1}} - e^{-\frac{TR}{T_1}} + 2e^{-\frac{TR - TE}{T_1}} \right] e^{-\frac{TE}{T_2}}, \quad (2)$$

where the scan parameters are echo time (TE), repetition time (TR) and inversion time (TI), and the signal evolution is evaluated for each pixel location. Generated example synthetic contrasts are shown in Figs. 4, 5, 6, and 7 for different inversion times. It can be clearly observed that the synthetic contrasts are different than the experimentally acquired ground truth contrasts shown in the first column. As the synthetically generated and true contrasts differ, in the next section we look at the possible unmodeled parameters and propose a deep learning framework to correct this mismatch.

2.2. Unmodeled parameters

There are many physical effects that are typically unmodeled when generating synthetic MR images from PD, T1, and T2 maps. Chief among them are RF-field (B1) inhomogeneity, diffusion and flow effects, and magnetization transfer (MT). In the following subsections, we briefly discuss their main impact on the observed image contrast.

2.2.1. B1 inhomogeneity

To flip the magnetization vectors from the steady state longitudinal direction, an RF pulse B₁ is applied whose strength is dependent upon the prescribed angle of rotation. The actual flip angle would not be same as the desired angle at all of the locations due to non-linear performance of the transmit system components responsible for B₁ magnetization. When this B₁ field inhomogeneity is not corrected, it can lead to errors in quantitative MRI parameter estimation [38,39]. If a signal with flip angle α_f is acquired, its location-dependent intensity would be [40].

$$I(r) = \sin(\alpha_f(r)) \cdot M_o, \quad (3)$$

where M_o is the longitudinal magnetization and $\alpha_f(r)$ is location dependent due to variation in B1 values based on location. Therefore, the outcome of this effect is a multiplicative term.

2.2.2. Diffusion

Diffusion weighting is based on the Brownian motion of water molecules in biological tissues, and it provides a powerful tool for investigating tissue microstructure and organization. The diffusion rate and direction will vary spatially due to differences in cellular membranes, tissue boundaries, etc. These can be probed by applying gradients separated by a particular delay. The net effect is a reduced magnetization which is dependent upon location and given as follows [41,42]:

$$M_{eff} = M_o \cdot e^{-\frac{\gamma^2}{2} \cdot \vec{b} \cdot \vec{D} \cdot \vec{b}} \quad (4)$$

where $\vec{D}$ is the diffusion coefficient that is unique for each spatial

location $\vec{r}$. $\vec{b}$ is a function of magnetic gradients ($\vec{G}(\tau)$), the time (t) for which gradients are applied, and the echo time:

$$\vec{b} = \gamma^2 \int_0^{TE} \left(\int_0^t \vec{G}(\tau) d\tau \right) d\hat{t}, \quad (5)$$

where γ is the gyromagnetic ratio. The multi-contrast sequence is likely to introduce some diffusion weighting due to gradients, and therefore the resulting images will have a spatial multiplicative factor.

2.2.3. Magnetization transfer

Magnetization transfer contrast is observed due to the transfer of protons in tissue from bound to free water molecules. When multiple saturation RF pulses are used, a new equilibrium is reached between the bound and free pools, leading to a change in contrast. The net effect is a reduced apparent longitudinal relaxation time T_1^* as compared to the normal T_1 relaxation time. Furthermore, due to MT, the total measurable MR signal at all voxels is decreased [43,44]. Formally, we can represent this as follows:

$$T_1^* = \frac{1}{R_1 + K_{for}}, \quad (6)$$

where $R_1 (= 1/T_1)$ is the longitudinal relaxation rate without the MT effects, and K_{for} is the magnetization transfer exchange rate between the bound and free pools of hydrogen nuclei. For an equilibrium longitudinal magnetization of M_0 , the effective magnetization due to MT effects is as follows:

$$M_{eff} = \frac{M_0}{1 + K_{for} \cdot T_1}. \quad (7)$$

As the multi-contrast sequence involves echo trains of RF pulses, the sequence will have inherent MT contrast. However, the sequence itself is not designed to estimate K_{for} . Nonetheless, it can be observed that the effect of reduced magnetization is a multiplicative term after generating the synthetic contrasts and a decrease in the T_1 value for different voxels.

2.2.4. Summary

In conclusion, the net effect of B_1 inhomogeneity, diffusion, and MT is a reduced total magnetization and T_1 relaxation time. Therefore each synthetic contrast will require a different multiplicative term to correct for unmodeled effects. The updated system model with multiplicative correction term is shown in Fig. 2. To compensate for the reduced T_1 value, we introduce an additional factor α into the synthetic contrast Eq. (2) as follows,

$$I = PD \left[1 - 2e^{-\frac{\alpha T}{T_1}} - e^{-\frac{T}{T_1}} + 2e^{-\frac{T - T_E}{T_1}} \right] e^{-\frac{TE}{T_2}} \quad (8)$$

Intuitively $\alpha < 1$ would help in countering the effect of reduced T_1 value and same is also observed in numerical experiments (Section 3.2). For this work, we have only included α compensation term when T_1 interacts with inversion time, as the experiments conducted in this work only explore the effects of variable T_1 .

2.3. Proposed approach

An overview of the proposed approach is shown in Fig. 2. A deep

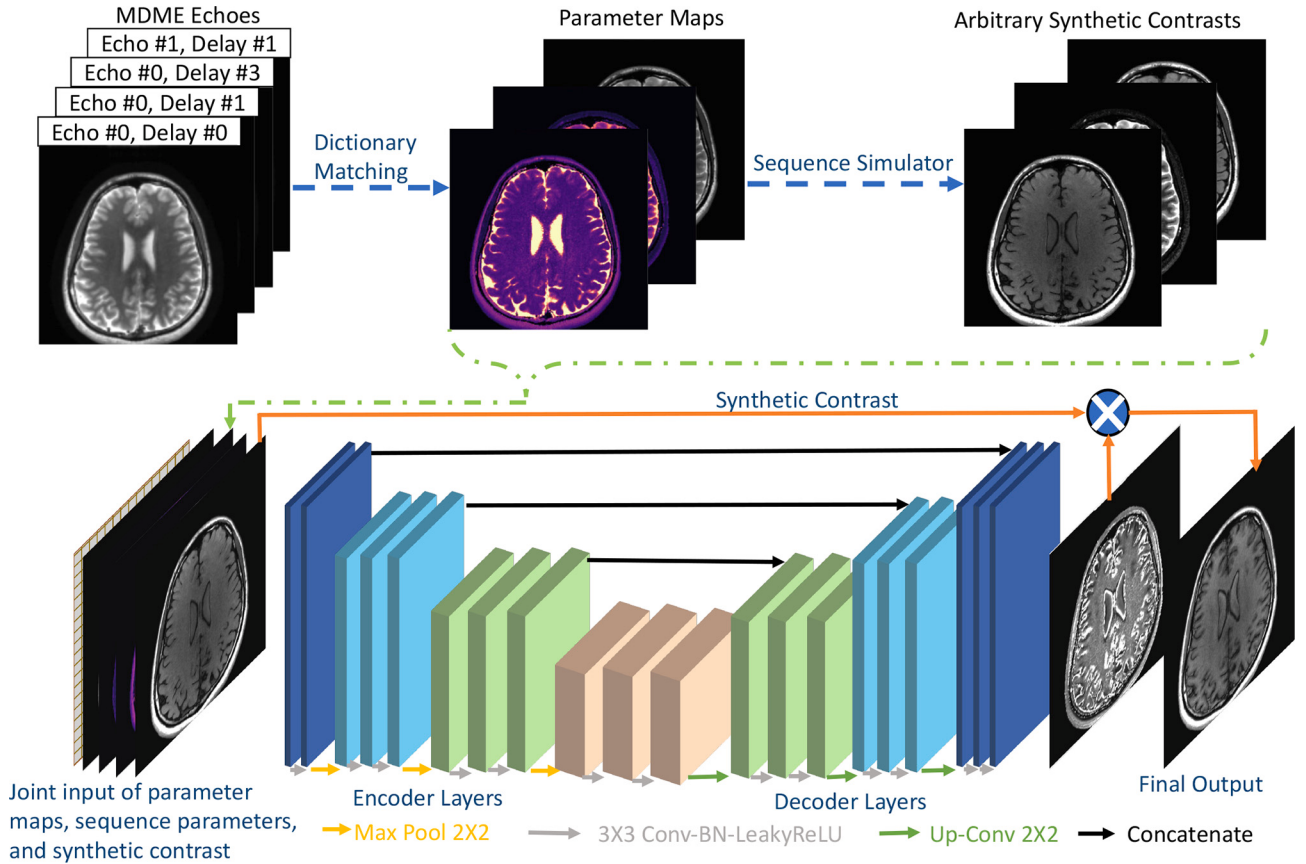


Fig. 2. Proposed contrast correction method. Parameter maps are estimated from the multi-contrast sequence (e.g. MDME echoes) and used to simulate arbitrary synthetic contrast images. The input to the neural network is the combination of parameter maps, sequence specifications, and the synthetic image contrast. Due to the inclusion of the three inputs, arbitrary contrasts can be corrected.

network takes as input the parameter maps (T_1 , T_2 , and PD) along with the synthetic MRI images and an additional channel with all values equal to normalized inversion time. Here the input synthetic image is generated using Eq. (8). The aim of the network is to generate the multiplicative term (M) and the final image output is $\hat{Y} = M \cdot I$, where I is obtained from Eq. (8). Next, we explain how the network architecture is selected along with the loss functions used to train it.

2.3.1. Network details

For the purpose of generating the multiplicative correction term from the synthetic MRI images and the parameter maps, we employ a Conditional GAN (cGAN) with multiple changes. cGANs have already been employed in other medical image synthesis and translation tasks [15,18–20]. GANs have two main components; a generator (G) and discriminator (D). The generator's task is to generate an output contrast given a particular set of inputs, whereas the job of discriminator is to distinguish whether the output image matches the distribution of reference contrast images or not.

For this work, we use a U-Net [45] as a generator, which has been extensively investigated for different medical imaging tasks [11,15]. To further inform the generator about the scan parameters, we include an additional channel with the same shape as the input image, but with a value at all pixels equal to the inversion time normalized by the maximum possible inversion time for the underlying anatomy (3000 ms in this work). This approach could be used for additional scan parameters.

For the discriminator, we initially explored a multi-layer convolutional network as described in [14]. While PatchGAN worked well for the application to ImageNet and other datasets, we found that it did not work suitably for our application, possibly due to the higher resolution of MR images. We decided to use a ResNet-18 model [46] with the number of output classes set to one as we found it to be more stable during training. The input channels were additionally reduced to two to work with the absolute value images and TI channel, making the discriminator aware of the TI used in the generator.

2.3.2. Loss functions

The network training is performed such that the final contrast-weighted images are as close to the reference scans as possible, such that the discriminator is not able to distinguish them. Therefore, both the generator and the discriminator are trained using an alternating optimization approach. The GAN minmax objective is formulated as follows:

$$\min_G \max_D \text{Obj}(D, G) = \mathbb{E}_{Y \sim p_{\text{data}}(Y)} [\log D(Y)] + \mathbb{E}_{Z \sim p_Z(z)} [\log(1 - D(I \cdot G(I, \theta, \mu)))] \quad (9)$$

where $\mathbb{E}(\cdot)$ is the expectation operator, Y are the reference contrast-weighted MR images, and z represents the input set of the generator consisting of synthetic MR images (I), parameter maps θ (e.g. T_1 , T_2 , and PD) and the scan parameters μ (i.e. TI in this work). The output $\hat{Y} = I \cdot G(I, \theta, \mu)$ tries to match the output distribution to be as close to the distribution of experimental contrast images.

The discriminator is trained with an adversarial loss whose objective is given in Eq. (9). The generator is trained with a combination of pixel-wise ℓ_2 -loss ($L_2 = \|Y - \hat{Y}\|_2^2$), adversarial loss (L_{adv}), and a feature preserving perceptual loss provided by the ImageNet pretrained VGG-19 network [47]. The perceptual loss is the ℓ_2 distance between the feature representations output by the last convolutional layer of the VGG network, given as follows [48],

$$\mathcal{L}_p(Y, \hat{Y}) = \frac{1}{H_{ij}W_{ij}} \sum_{x=1}^{W_{ij}} \sum_{z=1}^{H_{ij}} \|\psi_{ij}(Y) - \psi_{ij}(\hat{Y})\|_2^2 \quad (10)$$

where ψ_{ij} is the activation layer output from j^{th} max pooling layer and i^{th}

convolution with shape of H_{ij} , W_{ij} . Therefore the total generator loss after the combination of all these loss terms becomes,

$$\mathcal{L}_{\text{gen}} = L_2 + \lambda_1 L_{\text{adv}} + \lambda_2 L_p \quad (11)$$

where λ_1 and λ_2 are hyperparameters which are optimized heuristically over the validation set. Fig. 3 shows the overall training framework with the combination of different loss functions for the generator and discriminator networks.

2.3.3. Training optimization

To further improve performance on the proposed task, we make the following changes to the training algorithm from [14]. The input synthetic MRI images were normalized by the max value before being given as input to the network. Different learning rates are employed for the generator and discriminator as this was found to help training networks with vastly different architectures and training objectives. At each round, the discriminator was trained for more epochs compared to the generator (20 vs 10). Further, the overall training loop consisted of 20 epochs. We used weight clipping for the discriminator (0.01), as this is similar to the weight clipping proposed for the Wasserstein GAN [49]. Following that we also modified the discriminator loss by removing the logarithm from the last layer.

The discriminator was trained with a binary cross entropy loss combined with a sigmoid function for numerical stability, as suggested by [50]. We used the ADAM optimizer [51] with β set to [0.5, 0.999] instead of RMSprop as suggested in [49]. To find the exact numerical learning rates and hyperparameters, we tuned them sequentially. First, the learning rates for discriminator and generator were heuristically determined such that the training loss stabilized and there was no observed mode collapse. The exact numerical learning rates for discriminator and generator were 0.00005 and 0.0001, respectively. The first regularization parameter λ_1 was evaluated by linearly varying it while keeping $\lambda_2 = 0$. Once optimal λ_1 was tuned, we evaluated λ_2 by monitoring the adversarial loss. At optimal settings, the discriminator should output the same scores for both reference and generated images. The exact optimal numerical values found for λ_1 and λ_2 were 1.0 and 10.0, respectively.

2.4. Data acquisition

Healthy volunteers were recruited for brain MRI with institutional review board approval and informed consent. A total of 18 subjects were scanned on a Siemens 3 T Vida scanner (Siemens Healthineers, Erlangen, Germany) and 16-channel head coil, and 53 slices were acquired per subject. Of these, 14 subjects were used for training, one for validation, and three subjects for the test set. The central slices for each subject were used as these slices corresponded to predominant brain tissue. The overall scanning protocol consisted of the MDME sequence and several IR-FSE scans with a fixed TE of 11 ms, TR of 10,180 ms, and inversion time between 25 ms and 2500 ms such that the latent physical parameters can be learned by the network for all inversion times. The acquired multicoil k-space data were reconstructed using the BART toolbox [52] and quantitative PD, T1, and T2 maps were estimated using Python. The network was implemented using PyTorch and training was done using an NVIDIA GeForce RTX 3090 GPU with a batch size of 16. Training took approximately 3 h for 20 overall outer epochs, while inference took about 2 s per slice. Full scan parameters and relevant protocol details are provided in the next section.

2.5. Protocol settings

The MDME sequence is acquired with a slice thickness of 3 mm, and a field of view (FOV) of 22.8 cm. Further, the slices were acquired in the axial direction and 53 slices were acquired per subject. The TE values of 27 ms, 90 ms and TR value of 7680 ms were used to encode the T_2

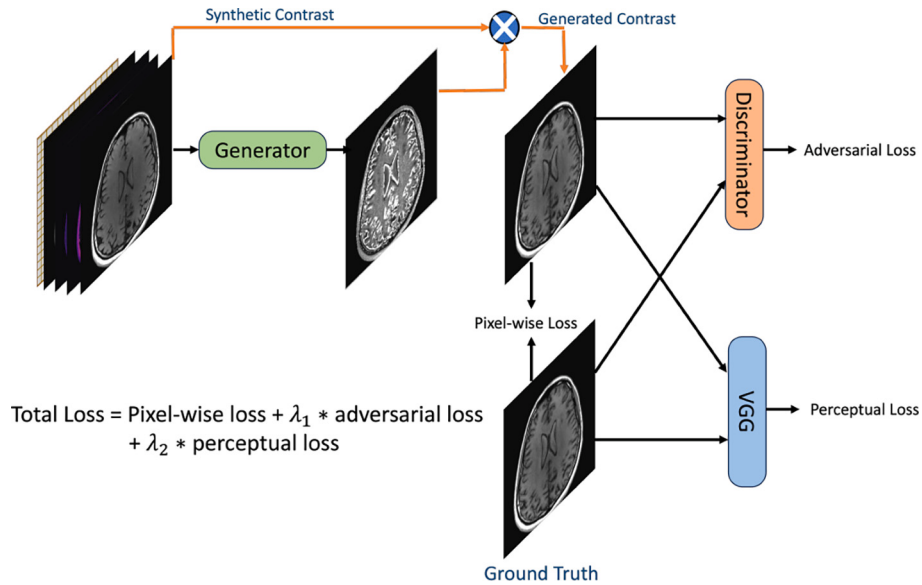


Fig. 3. The training of the proposed method is based on the conditional generative adversarial network training framework. The generator is trained using back-propagation on the total loss that incorporates pixel-wise L_2 loss, adversarial loss provided by the discriminator, and perceptual loss provided by the pre-trained VGG network. The discriminator is trained using the adversarial loss, while the VGG network is kept fixed.

information. Four delay times of 7562 ms, 3504 ms, 1041 ms, and 171 ms were used in the sequence to encode the T_1 information. The contrasts corresponding to different inversion times were acquired with the same FOV and resolution. The repetition and echo time are kept fixed at 10180 ms and 10 ms respectively. The original acquired k-space matrix size was 320×288 . The duration of MDME sequence was 8 mins whereas each individual contrast was about 3 mins long. The individual contrast data were acquired with an IR-FSE sequence which was further sub-sampled with an acceleration factor of 3.2.

2.6. Evaluation

To compare the effect of the multiplicative term, we trained three different models to correct the synthetic contrasts. In the residual model, the network output is added to the synthetic contrast (i.e $M + I$). In the direct model, the output of the network is the final contrast image. The proposed model multiplies the network output with the synthetic contrast. We compare both qualitative and quantitative image quality, where the latter is evaluated using both structural similarity index

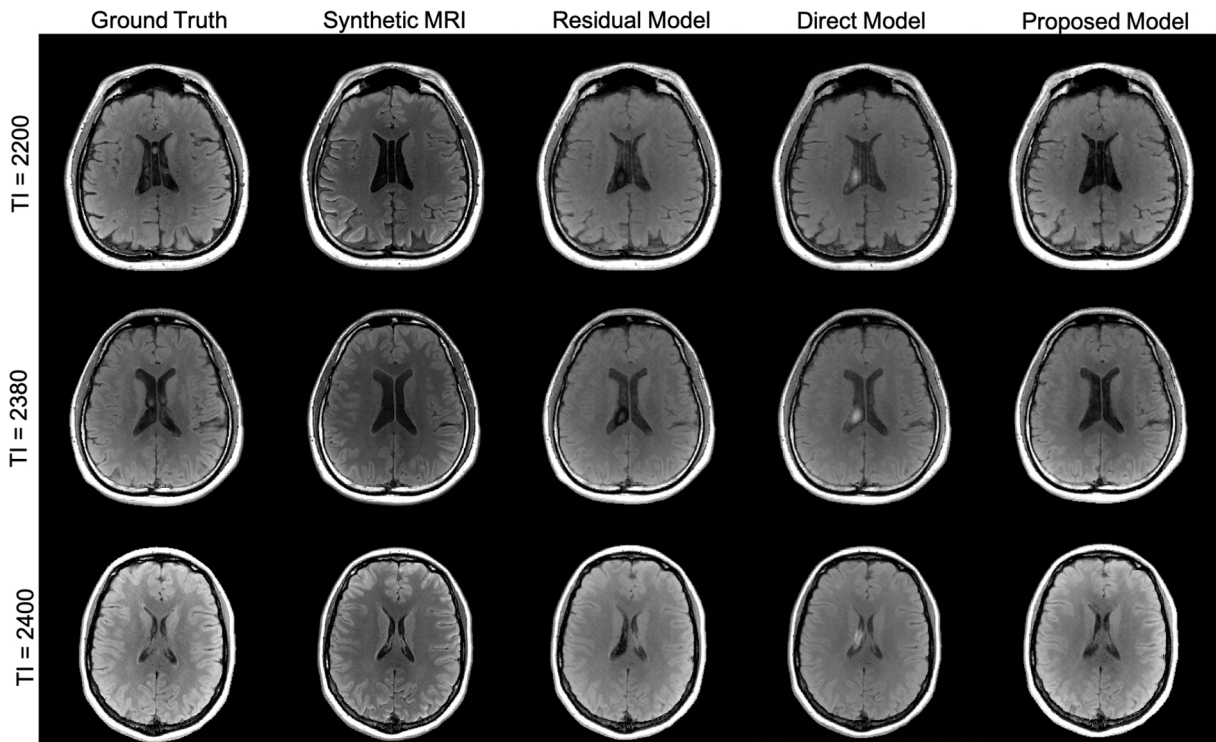


Fig. 4. Comparison of contrast correction using different models. The first column is for the ground truth data, the second column is for the synthetic MR image, third row is for the case of the residual model, the fourth row is for the direct contrast correction model and the last column is for the proposed multiplicative model. Each row corresponds to a different subject from the test set and different TI values (in ms) which are as mentioned in the Figure.

measure (SSIM) [53] and normalized root mean square error (NRMSE).

As numerical scores such as SSIM and NRMSE are known to be poor metrics, we additionally performed a blind reader study with a board-certified neurologist. The reader compared each of the proposed method and synthetic MRI images to the reference (conventional) contrast acquisition. The reader was blinded to the synthetic MRI and proposed method images, and they were randomly ordered for different generated contrasts. We chose 30 images from the test set displaying “Reference” on the left, and “A: Proposed” and “B: Synthetic MRI” on the right, where the order was randomly chosen. We asked the reader to rate the quality of the two methods with respect to the reference image on a scale of 1–5, 5 being best, for three categories: “Contrast”, “SNR”, and “Sharpness”. For each category, we performed a one-tailed paired *t*-test at a significance level of 0.05 with a null hypothesis that the score of the proposed method was not greater than the score of Synthetic MRI.

3. Results

3.1. Multiplicative model

Fig. 4 shows the ground truth, synthetic contrasts and the contrast correction results for the three different correction methods as mentioned in Section 2.6. Here the results are shown for TI values near 2000 ms to emphasize the nulling of CSF. Each row corresponds to a different test subject and a different TI value. The residual and direct models have error in the CSF region in the form of hyperintense and hypointense variation, likely due to flow and diffusion. The multiplicative model more faithfully matches the experimental contrast. Table 1 shows numerical error metrics for the different correction models. The proposed multiplicative model improves both SSIM and NRMSE metrics.

3.2. Comparison of effective T_1

Next, to compare how much effect the introduction of α factor (as introduced in Eq. (8)) has on the contrast correction we conducted the experiments over multiple values of α . Fig. 5 shows the contrast correction results for the different values of α . The first column is for the ground truth contrasts and after that, the corrected contrasts are shown sequentially for 5 different α values of 1.4, 1.2, 1.0, 0.8, and 0.6. Four rows of results corresponding to different TI values are shown. Due to magnetization transfer, the effective T_1 value of tissues decreases, and adding a α value smaller than 1, helps in alleviating that effect. A similar trend is observed in the results where the higher α value corrected contrasts still have some signal left in the CSF region. Moreover, there are other contrast mismatches in the white and the gray matter of the brain. It can be clearly observed in the CSF region that with the decrease of α factor the residual signal in the CSF region is also reduced. Table 2 shows the numerical error metrics for models with different α values and it can be observed that the metric first improves with increasing α and then start to decrease. Therefore we chose $\alpha = 0.8$ and held it fixed for all remaining results.

Table 1

Performance comparison of different methods for contrast translation on the validation set. The mean value along with the standard deviation of all the error metrics is provided. The proposed method gives the best performance among all the possible correction techniques.

Error metric	SSIM	NRMSE
Synthetic MRI	0.778 $\pm$ 0.073	0.456 $\pm$ 0.117
Direct Model	0.806 $\pm$ 0.053	0.419 $\pm$ 0.119
Residual Model	0.819 $\pm$ 0.059	0.317 $\pm$ 0.098
Proposed Model	0.846 $\pm$ 0.062	0.303 $\pm$ 0.132

3.3. Results over the full TI range

Having shown that the multiplicative proposed model with $\alpha = 0.8$ provides best qualitative and quantitative results, next we show the contrast correction results over the full range of inverse times. In that regard, Fig. 6 shows the contrast correction results for 5 different TI values ranging from 100 ms to 2430 ms over the test set. The first row shows the results for the ground truth, the second row shows synthetic MRI contrasts, and the third row shows the corrected contrasts according to the proposed method. The last row shows the direct contrast correction [11]. It can be clearly seen for TI = 1330 that there is a residual signal in the CSF region for the synthetic MRI contrast whereas the proposed method correctly nulls it. For TI values of 1830 and 2430 ms, the synthetic contrast correctly nulls the CSF region but exhibits contrast mismatch between the gray and white matter, which is corrected in the proposed method. Direct contrast synthesis shows excellent synthetic contrast, but it is only able to produce contrasts that were explicitly collected at scan time and hence available in the training set. Therefore, the higher TI values that were not in the training set cannot be synthesized. We are not aware of another deep learning-based method that can correct contrasts at arbitrary inversion times that were not seen during training.

3.4. Impact of training objective

Fig. 7 shows the experimental contrasts, synthetic contrasts, and corrected contrasts for three arbitrarily chosen inversion times in the test set for different combinations of loss functions. Synthetic MRI column represent the image generated through Bloch equations (c.f. 2). The third and fourth columns show the result with L1-loss and with L1 + GAN loss, respectively. The last column shows the results with L1, adversarial and perceptual loss combined. The inclusion of perceptual loss qualitatively retains sharper features. While the synthetic contrasts show high resolution, they distort the contrast due to the simplified signal model. The proposed method improves the contrast while preserving most fine features.

Table 3 shows the numerical error metrics for the different training objectives. The UNET gives the best SSIM and NRMSE metrics. This is an expected result as the UNET is only trained with a L2-loss (which closely corresponds to NRMSE), and it has been shown to be proportional to these loss metrics, even though these low loss value come at a cost of blurring and loss of resolution in final images. Therefore we also included the perceptual loss value which better captures the visual quality. The proposed method gives the best perceptual loss metric. For lower TI values, both UNET and GAN based methods can correct the contrast but substantially reduce the sharpness of the image, whereas the addition of perceptual loss retains the fine features as well as corrects the contrast mismatch. For the higher value of TI = 1330 as shown in the third row, the UNET and GAN based methods were not able to remove the signal in from the CSF region as well as there is a leaking hazing effect due to the low resolution of these methods.

3.5. Reader study scores

Fig. 8 shows the bar plot results of the clinical reader study. For all three evaluation metrics contrast, SNR, and sharpness, the proposed method has a higher average score than the synthetic MR images. We found that contrast and SNR results are statistically significant with a *P*-value < 0.01. The statistical parameters of the reader study are summarized in Table 4.

4. Discussion

Synthetic MRI has the potential to greatly improve scan efficiency and reduce scan times. Furthermore, the ability to retrospectively synthesize new contrasts opens opportunities for adding new sequences in

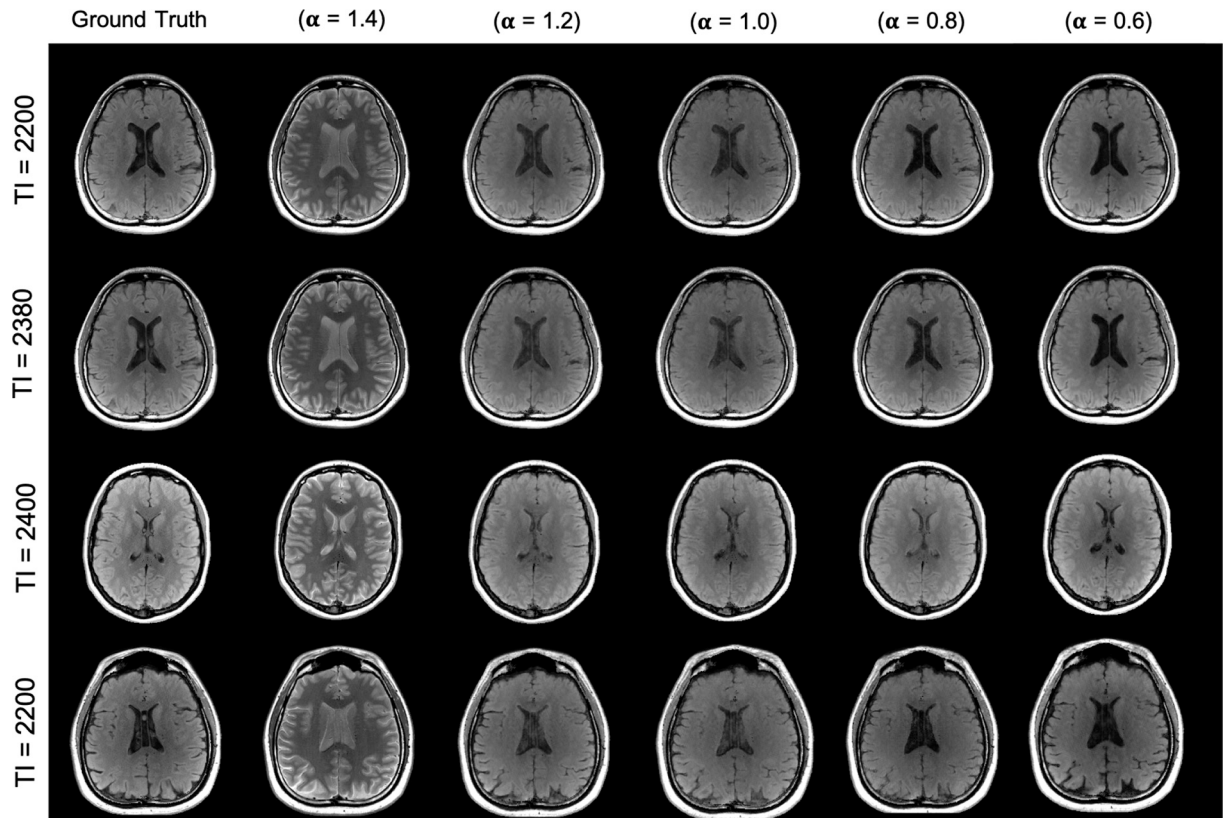


Fig. 5. Comparison of contrast correction for different α values as used in the Eq. (8). The contrasts are shown for five different α values of 1.4, 1.2, 1.0, 0.8 and 0.6. The first two rows are for the same subject whereas the last two rows are for different subjects.

Table 2

Performance comparison of different methods for contrast translation on the validation set. The mean value along with the standard deviation of all the error metrics is provided. $\alpha = 0.8$ provides the best performance among all the experimental α values.

Error metric	SSIM	NRMSE
$\alpha = 0.6$	0.821 ± 0.069	0.332 ± 0.102
$\alpha = 0.8$	0.846 ± 0.062	0.303 ± 0.132
$\alpha = 1.0$	0.821 ± 0.061	0.370 ± 0.169
$\alpha = 1.2$	0.820 ± 0.064	0.382 ± 0.179
$\alpha = 1.4$	0.813 ± 0.063	0.401 ± 0.208

silico as well as future “false contrasts” that highlight particular features. By using an approach that incorporates important MRI sequence parameters along with synthetic MRI images, this work generates arbitrary synthetic contrasts that more faithfully match experimental scans. Other deep learning based methods directly translate from a base multi-contrast image to the target contrast and doing so loses the ability to retrospectively generate arbitrary contrasts. To show how our method is different, we plotted comparison results in Fig. 6, where the last row are the results that we would get if the model only learns to translate from one image to another and those TI values are already present in the training set. Whereas, the last three columns show the contrast correction for three different inversion times which were not present in the training set.

3D Synthetic MRI is an exciting emerging approach, and we did not explore it in this work. We emphasize that our methodology is flexible in that the input images can be generated with different multi-contrast acquisitions, although the data collection and training would need to be repeated for that specific acquisition. This is the case for all existing deep learning-based synthetic MRI methods and is not unique to ours. Therefore, our methodological framework was tested with the sequences

that were available to us on our system, though future work should explore the use of different sequences (both 2D and 3D), for example, 2D and 3D MR Fingerprinting [3,54] and 3D-QALAS [55].

For model comparison and optimal α evaluation, we only show the results (Figs. 4 and 5) for higher TI values as these are the contrasts that are most susceptible to mismatch in the CSF due to unmodeled effects. We showed that naively applying deep learning contrast correction without regard to the nature of the unmodeled effects can lead to poor contrast generation. Our ablation study of the loss function does show the importance of incorporating both pixelwise error as well as perceptual error in the form of adversarial and perceptual losses. Analogous to prior work [11], we did observe that the vanilla UNET with pixelwise L_2 loss gives a better SSIM and NRMSE metric as shown in Table 3, even though those images are blurry and may introduce artifacts as shown in the Fig. 7. However, from visual inspection, it can be concluded that the GAN with perceptual loss gives better quality images that are free of smoothing around boundaries and are more accurate. This is further corroborated by the lower perceptual loss metric for the proposed method. Specifically, there is still some signal in the CSF region for the UNET and GAN output due to partial voluming effects, which the inclusion of perceptual loss helps in alleviating. This shows that quantitative metrics alone are not sufficient to judge image quality, which is a known issue in the literature.

To further corroborate the hypothesis that the proposed method improves image quality, results from a blind reader study were included. These results show with statistical significance that the proposed method has improved both contrast and SNR. There are improvements in the mean sharpness, however, those were not statistically significant. It can be clearly observed from that the quantitative measures of error i. e. NRMSE and SSIM doesn't correlate with the quality, which may also be due to image alignment between scans. The same observation has been made by multiple previous works [10,11].

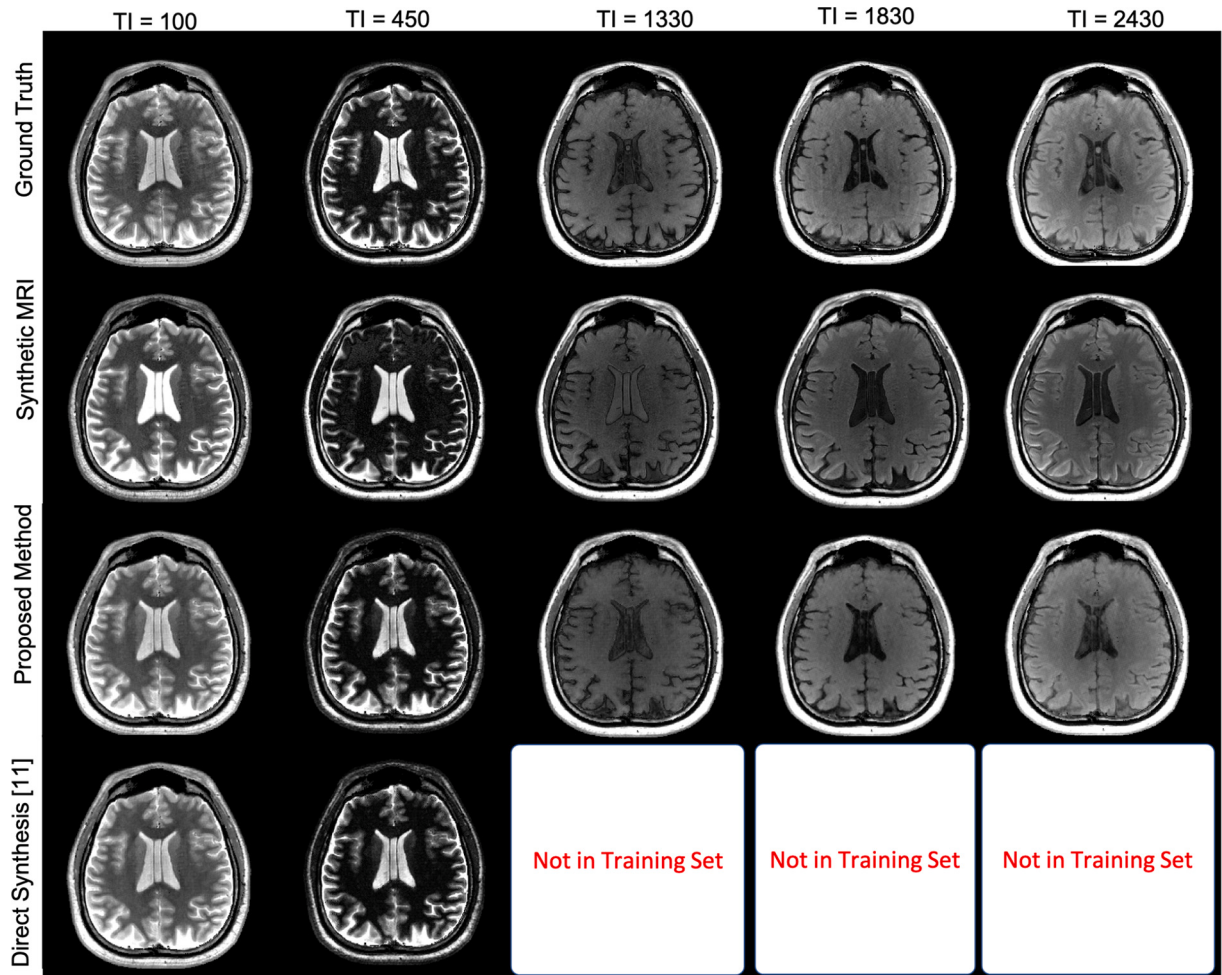


Fig. 6. Example contrast corrections for 5 different inversion times. The first row is the ground truth contrasts, the second row is for the synthetic MRI contrasts and the third row is for the corrected contrasts according to the proposed method. To do a fair comparison, all the images are from a single test subject.

With MR images, there is always a possibility of motion corruption, noise in the reconstruction, ringing artifacts, image registration across different contrasts, and other artifacts. Even though the introduction of perceptual loss helps in retaining fine features, it depends upon the dataset on which the VGG network is trained and it is also known to introduce checkerboard patterns in the output images [56]. The hyperparameter tuning and training optimization was performed to find a stable training landscape between the generator and discriminator while also qualitatively producing high-quality images over the validation set. Nonetheless, there is still a risk of hallucinations or artifacts being introduced by complicated deep learning based processing pipelines.

Our approach is able to correct synthetic MRI at arbitrary contrasts. However, as it is a deep learning based method it is still susceptible to overfitting to the training set. To understand this effect, we plot the inversion recovery curve across the full TI range for three different regions of interest (ROI) in Fig. 9. The plot shows the inversion recovery for the contrasts generated using synthetic MRI and the proposed method for a single slice of one of the test subjects. The ROIs correspond to white matter, gray matter, and CSF were manually selected, and the mean signal values were taken over the ROIs. Additionally, the experimentally acquired data points are shown using as points on the plot for the corresponding tissue types for the limited contrasts that could be captured per volunteer due to scan time limitations. The synthetically generated contrasts curve follows the exponential recovery curve based on (2), and deviate from the experimental points. It can be observed that the proposed model output curve follows the experimental data points

better at most of the inversion times, though there is clear evidence of overfitting at low TI values, likely due to the smaller number of TI contrasts collected around that region. Therefore, even though the proposed approach could correct arbitrary contrasts, it is critical to obtain sufficient training data across many subjects. Notably, it is not necessary to scan each subject with every training point. Nonetheless, the results show promise in a generalizable method for contrast generation and correction for arbitrary scan parameters. Furthermore, Table 5 shows the NRMSE for the signal value of different tissues as compared to the experimental data over the ROIs. The proposed method has lower error for all tissue types as compared with synthetic MRI. This observation indicates that the proposed method better fits the experimental data in terms of the actual signal value at different inversion times.

There are many choices for the discriminator network, which could affect the training. We chose the ResNet-18 in this work after internal experimentation with several architectures including PatchGAN, and found that the ResNet-18 is more stable during training. It is possible that other multi-layer convolutional network architectures could also be used, though we found that the ResNet-18 based discriminator had good classification efficiency by the end of training.

One limitation of this work is that our training data was acquired by only varying the TI values and did not include TE and TR variation. We therefore do not know how much training data are required to train a network for full range of TI, TE, and TR values and how difficult that would be. As the method is already prone to overfitting, this limitation is important to consider. In addition, as more experimental scans are acquired, subject motion becomes an issue. Coregistering images from all

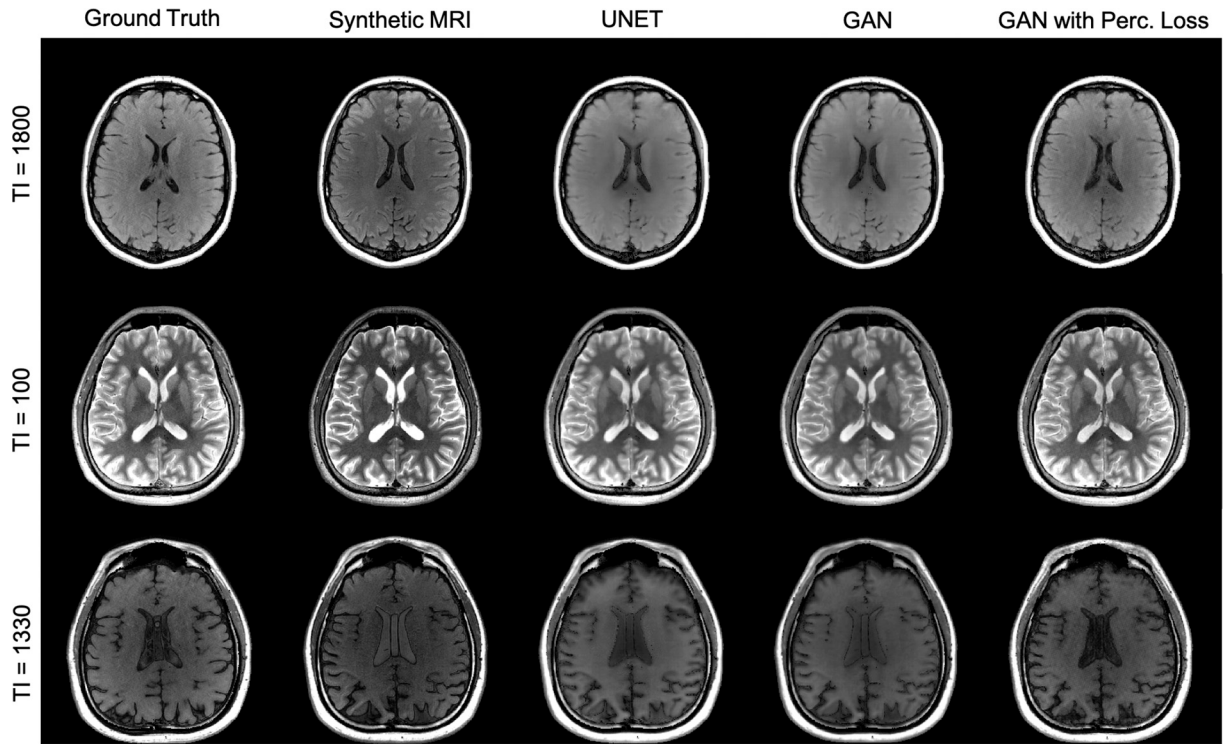


Fig. 7. Example contrast corrections for 3 different inversion times. Each row corresponds to a different test subject. Third to fifth column corresponds to a different method of neural network training.

Table 3

Performance comparison of different training objectives for contrast translation on the validation set. The mean value along with the standard deviation of all the error metrics is provided.

Error metric	SSIM	NRMSE	Perc. Loss
UNET	0.856 $\pm$ 0.06	0.283 $\pm$ 0.099	1.091 $\pm$ 0.279
+ GAN	0.853 $\pm$ 0.006	0.326 $\pm$ 0.119	1.093 $\pm$ 0.273
+ Perc. Loss	0.846 $\pm$ 0.062	0.303 $\pm$ 0.132	1.011 $\pm$ 0.263

Table 4

Statistical results from the reader study.

	Contrast	SNR	Sharpness
Synthetic MRI	4.13 $\pm$ 0.43	4.6 $\pm$ 0.49	4.43 $\pm$ 0.5
Proposed	4.7 $\pm$ 0.59	4.87 $\pm$ 0.34	4.63 $\pm$ 0.48
T-statistic	3.319	2.804	1.533
P-value	0.0024	0.0089	0.1360

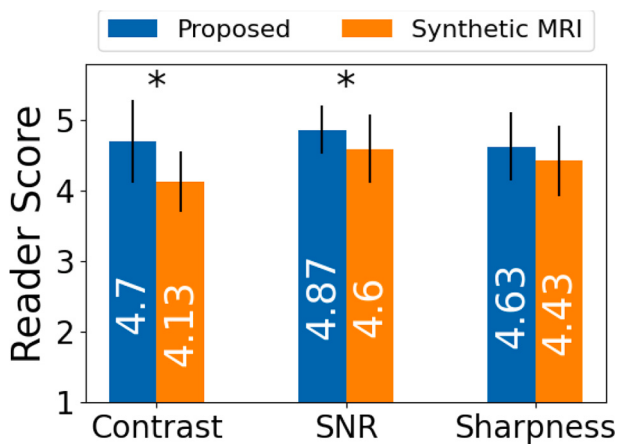


Fig. 8. Bar plot results of clinical reader study. For all metrics, the proposed method has a higher average score. The error bars represent the standard deviation and * above a barplot represents statistical significance with p -value under 0.01.

the different acquired contrasts may be necessary.

Another limitation of this work is the acquired dataset only includes healthy subjects and therefore does not contain pathology. We

hypothesize that the generality of the approach could be extended to generating contrasts with pathology. However, in this present study we were not able to investigate this. It is important to explore this in clinical validation as a next step. Finally, though we show that a multiplicative factor is a more natural fit for contrast correction, it is still a heuristic that does not fully account for unmodeled effects, and therefore will need further exploration.

5. Conclusion

In this paper, we proposed a novel deep learning framework to correct the contrasts generated using synthetic MRI. The proposed method incorporated unmodeled physical effects through the use of a multiplicative correction network and an additional effective T1 correction. Our method generates the correction image as a scan parameter-informed image-to-image translation using a conditional GAN trained along with a perceptual loss. To inform the network about different contrasts, the quantitative maps, synthetic MR images, and scan parameters as an additional channel are given as input during the training process. The results show the improved performance of the proposed multiplicative model over other deep learning based correction methods. Furthermore, we were able to correct contrasts that were not present in the training set for different subjects. The results show improved performance in terms of quantitative error metrics as well as qualitatively. A possible extension of this work is to extend the contrast correction for arbitrary echo times and repetition times, which will

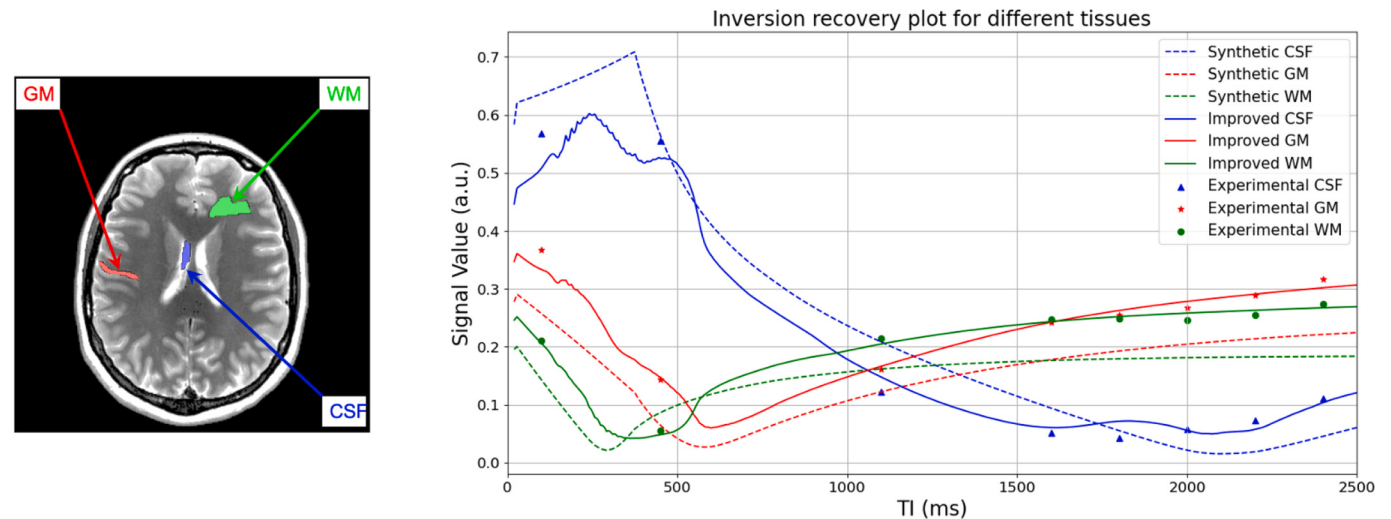


Fig. 9. Inversion recovery plot for three different tissue types for the ground truth, synthetic MRI, and proposed model improved images. CSF: cerebrospinal fluid; WM: white matter; GM: gray matter. The left column shows the different regions selected corresponding to different tissues.

Table 5
NRMSE comparison of synthetic MRI method and the proposed method against the experimental contrast tissue signal values over the ROIs. The proposed method has lower numerical error values for three tissue regions.

Method	Synthetic	Proposed
CSF	0.1843	0.0972
GM	0.2878	0.1162
WM	0.2917	0.1166

likely require more training data to ensure that the network learns the underlying unmodeled effects. Another potential use of the proposed method could be to estimate the unmodeled physics in order to update the quantitative parameter maps.

Disclosures

L.F. has received fees for consultancy and/or advisory board participation from Genentech, Roche, Novartis, Bristol Myers Squibb, EMD Serono, Sanofi, Horizon Therapeutics and TG Therapeutics; has received honorarium for participation in educational programs from Medscape, Inc. and the MS Association of America; has received speaking fees from the MS Association of America, EMD Serono, and Sanofi; has received program sponsorship from EMD Serono and grant support from NIH/NINDS, PCORI, Genentech, and EMD Serono through her institution.

CRedit authorship contribution statement

Sidharth Kumar: Conceptualization, Data curation, Formal analysis, Investigation, Methodology, Software, Visualization, Writing – original draft, Writing – review & editing. **Hamidreza Saber:** Validation. **Odelin Charron:** Conceptualization, Investigation, Supervision, Validation. **Leorah Freeman:** Conceptualization, Supervision. **Jonathan I. Tamir:** Conceptualization, Data curation, Formal analysis, Funding acquisition, Investigation, Supervision, Validation, Writing – original draft, Writing – review & editing.

Acknowledgements

This work was supported by NSF CCF-2239687 (CAREER), NSF IFML 2019844, and NIH U24EB029240.

References

[1] Simon J, Li D, Traboulsee A, Coyle P, Arnold D, Barkhof F, et al. Standardized MR imaging protocol for multiple sclerosis: consortium of ms centers consensus guidelines. *Am J Neuroradiol* 2006;27(2):455–61.

[2] Slipsager JM, Glimberg SL, Søgaard J, Paulsen RR, Johannessen HH, Martens PC, et al. Quantifying the financial savings of motion correction in brain MRI: a model-based estimate of the costs arising from patient head motion and potential savings from implementation of motion correction. *J Magn Reson Imaging* 2020;52(3):731–8.

[3] Ma D, Gulani V, Seiberlich N, Liu K, Sunshine JL, Duerk JL, et al. Magnetic resonance fingerprinting. *Nature* 2013;495(7440):187–92.

[4] Wamntjes J, Dahlqvist O, Lundberg P. Novel method for rapid, simultaneous t1, t* 2, and proton density quantification. *Magn Reson Med* 2007;57(3):528–37.

[5] Skare S, Sprenger T, Norbeck O, Rydén H, Blomberg L, Avventi E, et al. A 1-minute full brain MR exam using a multicontrast epi sequence. *Magn Reson Med* 2018;79(6):3045–54.

[6] Wang F, Dong Z, Reese TG, Bilgic B, Katherine Manhard M, Chen J, et al. Echo planar time-resolved imaging (epti). *Magn Reson Med* 2019;81(6):3599–615.

[7] Wamntjes J, Leinhard OD, West J, Lundberg P. Rapid magnetic resonance quantification on the brain: optimization for clinical usage. *Magn Reson Med* 2008;60(2):320–9.

[8] Cloos MA, Knoll F, Zhao T, Block KT, Bruno M, Wiggins GC, et al. Multiparametric imaging with heterogeneous radiofrequency fields. *Nat Commun* 2016;7(1):12445.

[9] Vargas M, Boto J, Delatre B. Synthetic MR imaging sequence in daily clinical practice. *Am J Neuroradiol* 2016;37(10):E68–9. <https://doi.org/10.3174/ajnr.A4895>. URL: <http://www.ajnr.org/content/37/10/E68>.

[10] Wang G, Gong E, Banerjee S, Martin D, Tong E, Choi J, et al. Synthesize high-quality multi-contrast magnetic resonance imaging from multi-echo acquisition using multi-task deep generative model. *IEEE Trans Med Imaging* 2020;39(10):3089–99.

[11] Wang K, Doneva M, Meineke J, Amthor T, Karasan E, Tan F, et al. High-fidelity direct contrast synthesis from magnetic resonance fingerprinting. *Magn Reson Med* 2023;90:2116–29.

[12] Hagiwara A, Otsuka Y, Hori M, Tachibana Y, Yokoyama K, Fujita S, et al. Improving the quality of synthetic flair images with deep learning using a conditional generative adversarial network for pixel-by-pixel image translation. *Am J Neuroradiol* 2019;40(2):224–30.

[13] Virtue P, Tamir JI, Doneva M, Yu SX, Lustig M. Direct contrast synthesis for magnetic resonance fingerprinting. In: *Proceedings of the 2018 ISMRM Workshop on Machine Learning*; 2018. p. 676.

[14] Isola P, Zhu J-Y, Zhou T, Efros AA. Image-to-image translation with conditional adversarial networks. In: *Proceedings of the IEEE conference on computer vision and pattern recognition*; 2017. p. 1125–34.

[15] Pasumarthi S, Tamir JI, Christensen S, Zaharchuk G, Zhang T, Gong E. A generic deep learning model for reduced gadolinium dose in contrast-enhanced brain MRI. *Magn Reson Med* 2021;86(3):1687–700.

[16] Dar SU, Yurt M, Karacan L, Erdem A, Erdem E, Çukur T. Image synthesis in multi-contrast MRI with conditional generative adversarial networks. *IEEE Trans Med Imaging* 2019;38(10):2375–88.

[17] Zhu J-Y, Park T, Isola P, Efros AA. Unpaired image-to-image translation using cycle-consistent adversarial networks. In: *Proceedings of the IEEE international conference on computer vision*; 2017. p. 2223–32.

[18] Armanious K, Jiang C, Fischer M, Küstner T, Hepp T, Nikolaou K, et al. Medgan: medical image translation using gans. *Comput Med Imaging Graph* 2020;79:101684.

- [19] Manaswi NK. Generative adversarial networks with industrial use cases: Learning how to build GAN applications for retail. Healthcare, Telecom, Media, Education, and HRTech, BPB Publications; 2020.
- [20] Xue Y, Xu T, Zhang H, Long LR, Huang X. Segan: adversarial network with multi-scale l1 loss for medical image segmentation. *Neuroinformatics* 2018;16(3): 383–92.
- [21] Yurt M, Dar SU, Erdem A, Erdem E, Oguz KK, Çukur T. Mustgan: multi-stream generative adversarial networks for MR image synthesis. *Med Image Anal* 2021;70: 101944.
- [22] Levac B, Kumar S, Kardoni S, Tamir JI. FSE compensated motion correction for MRI using data driven methods. In: *Proceedings of the Medical Image Computing and Computer Assisted Intervention – MICCAI*; 2022. p. 707–16.
- [23] Sharma A, Hamarneh G. Missing MRI pulse sequence synthesis using multi-modal generative adversarial network. *IEEE Trans Med Imaging* 2019;39(4):1170–83.
- [24] Zhou T, Fu H, Chen G, Shen J, Shao L. Hi-net: hybrid-fusion network for multi-modal MR image synthesis. *IEEE Trans Med Imaging* 2020;39(9):2772–81.
- [25] Yu B, Zhou L, Wang L, Shi Y, Frapp J, Bourgeat P. Ea-gans: edge-aware generative adversarial networks for cross-modality MR image synthesis. *IEEE Trans Med Imaging* 2019;38(7):1750–62.
- [26] Kwon G, Han C, Kim D-S. Generation of 3d brain MRI using auto-encoding generative adversarial networks. In: *International conference on medical image computing and computer-assisted intervention*. Springer; 2019. p. 118–26.
- [27] Abramian D, Eklund A. Refacing: reconstructing anonymized facial features using gans. In: *2019 IEEE 16th international symposium on biomedical imaging (ISBI 2019)*. IEEE; 2019. p. 1104–8.
- [28] Dai X, Lei Y, Fu Y, Curran WJ, Liu T, Mao H, et al. Multimodal MRI synthesis using unified generative adversarial networks. *Med Phys* 2020;47(12):6343–54.
- [29] Li H, Paetzold JC, Sekuboyina A, Kofler F, Zhang J, Kirschke JS, et al. Diamondgan: Unified multi-modal generative adversarial networks for MRI sequences synthesis. In: *International conference on medical image computing and computer-assisted intervention*. Springer; 2019. p. 795–803.
- [30] Liu F, McMillan A. MR image synthesis using a deep learning based data-driven approach. In: *Proc. 26th Annu. Meeting (ISMRM)*; 2018. p. 3490.
- [31] Gong E, Banerjee S, Pauly J, Zaharchuk G. Improved synthetic MRI from multi-echo MRI using deep learning. In: *Proceedings of 26th Annual Meeting of ISMRM*, Paris, France. Vol. 2795; 2018.
- [32] Ji S, Yang D, Lee J, Choi SH, Kim H, Kang KM. Synthetic MRI: technologies and applications in neuroradiology. *J Magn Reson Imaging* 2022;55(4):1013–25.
- [33] Kleinlog JP, Mandija S, D'Agata F, Liu H, van der Heide O, Koktas B, et al. Synthetic MRI with magnetic resonance spin tomography in time-domain (mr-stat): results from a prospective cross-sectional clinical trial. *J Magn Reson Imaging* 2023;57:14451–61.
- [34] Tanenbaum LN, Tsiouris AJ, Johnson AN, Naidich TP, DeLano MC, Melhem ER, et al. Synthetic MRI for clinical neuroimaging: results of the magnetic resonance image compilation (magic) prospective, multicenter, multireader trial. *Am J Neuroradiol* 2017;38(6):1103–10.
- [35] Ryu KH, Baek HJ, Moon JI, Choi BH, Park SE, Ha JY, et al. Initial clinical experience of synthetic MRI as a routine neuroimaging protocol in daily practice: a single-center study. *J Neuroradiol* 2020;47(2):151–60.
- [36] Weigel M. Extended phase graphs: dephasing, rf pulses, and echoes-pure and simple. *J Magn Reson Imaging* 2015;41(2):266–95.
- [37] Kingsley PB, Ogg RJ, Reddick WE, Steen RG. Correction of errors caused by imperfect inversion pulses in MR imaging measurement of t1 relaxation times. *Magn Reson Imaging* 1998;16(9):1049–55.
- [38] Sengupta A, Gupta RK, Singh A. Evaluation of b1 inhomogeneity effect on DCE-MRI data analysis of brain tumor patients at 3t. *J Transl Med* 2017;15(1):1–13.
- [39] Gonçalves FG, Serai SD, Zucconi G. Synthetic brain MRI: review of current concepts and future directions. *Top Magn Reson Imaging* 2018;27(6):387–93.
- [40] Cunningham CH, Pauly JM, Nayak KS. Saturated double-angle method for rapid b1 + mapping, magnetic resonance in medicine: an official journal of the international society for. *Magn Reson Med* 2006;55(6):1326–33.
- [41] Le Bihan D, Mangin J-F, Poupon C, Clark CA, Pappata S, Molko N, et al. Diffusion tensor imaging: concepts and applications, journal of magnetic resonance imaging: an official journal of the international society for. *Magn Reson Med* 2001;13(4): 534–46.
- [42] McNab JA. The bloch-Torrey equation & diffusion imaging (dwi, dti, q-space imaging). In: *Proceedings of the 21th annual meeting of ISMRM*; 2013.
- [43] Wolff SD, Balaban RS. Magnetization transfer contrast (mtc) and tissue water proton relaxation in vivo. *Magn Reson Med* 1989;10(1):135–44.
- [44] Boer R. Magnetization transfer contrast part 1: MR physics. *Medicamundi* 1995;40 (2):64–73.
- [45] Ronneberger O, Fischer P, Brox T. U-net: Convolutional networks for biomedical image segmentation. In: *International conference on medical image computing and computer-assisted intervention*. Springer; 2015. p. 234–41.
- [46] He K, Zhang X, Ren S, Sun J. Deep residual learning for image recognition. In: *Proceedings of the IEEE conference on computer vision and pattern recognition*; 2016. p. 770–8.
- [47] Simonyan K, Zisserman A. Very deep convolutional networks for large-scale image recognition. *International Conference on Learning representations*. 2015.
- [48] Johnson J, Alahi A, Fei-Fei L. Perceptual losses for real-time style transfer and super-resolution. In: *European conference on computer vision*. Springer; 2016. p. 694–711.
- [49] Arjovsky M, Chintala S, Bottou L. Wasserstein generative adversarial networks, in: *International conference on machine learning*. PMLR 2017:214–23.
- [50] P. Contributors. BCEWITHLOGITLOSS. URL, <https://pytorch.org/docs/stable/generated/torch.nn.BCEWithLogitsLoss.html>; 2023.
- [51] Kingma DP, Ba J. Adam: A method for stochastic optimization. *International Conference on Learning representations*. 2015.
- [52] Uecker M, Ong F, Tamir JI, Bahri D, Virtue P, Cheng JY, et al. Berkeley advanced reconstruction toolbox. In: *Proc. Intl. Soc. Mag. Reson. Med* 23. Vol. 23; 2015. p. 2486.
- [53] Wang Z, Bovik AC, Sheikh HR, Simoncelli EP. Image quality assessment: from error visibility to structural similarity. *IEEE Trans Image Process* 2004;13(4):600–12.
- [54] Ma D, Jiang Y, Chen Y, McGivney D, Mehta B, Gulani V, et al. Fast 3d magnetic resonance fingerprinting for a whole-brain coverage. *Magn Reson Med* 2018;79(4): 2190–7.
- [55] Kvernby S, Warntjes MJB, Haraldsson H, Carlhäll C-J, Engvall J, Ebbers T. Simultaneous three-dimensional myocardial t1 and t2 mapping in one breath hold with 3d-qals. *J Cardiovasc Magn Reson* 2014;16(1):1–14.
- [56] Wang K, Tamir JI, De Goyeneche A, Wollner U, Brada R, Yu S, et al. High fidelity deep learning-based MRI reconstruction with instance-wise discriminative feature matching loss. *Magn Reson Med* 2022;88(1):476–91.