



This MICCAI paper is the Open Access version, provided by the MICCAI Society. It is identical to the accepted version, except for the format and this watermark; the final published version is available on SpringerLink.

MDPG: Multi-domain Diffusion Prior Guidance for MRI Reconstruction

Lingtong Zhang, Mengdie Song, Xiaohan Hao, Huayu Mai, Bensheng Qiu^(✉)

School of Information Science and Technology,
University of Science and Technology of China, Hefei, China
{zhanglingtong, smengdie, hxx045, mai556}@mail.ustc.edu.cn
bqiu@ustc.edu.cn

Abstract. Magnetic Resonance Imaging (MRI) reconstruction is essential in medical diagnostics. As the latest generative models, diffusion models (DMs) have struggled to produce high-fidelity images due to their stochastic nature in image domains. Latent diffusion models (LDMs) yield both compact and detailed prior knowledge in latent domains, which could effectively guide the model towards more effective learning of the original data distribution. Inspired by this, we propose Multi-domain Diffusion Prior Guidance (MDPG) provided by pre-trained LDMs to enhance data consistency in MRI reconstruction tasks. Specifically, we first construct a Visual-Mamba-based backbone, which enables efficient encoding and reconstruction of under-sampled images. Then pre-trained LDMs are integrated to provide conditional priors in both latent and image domains. A novel Latent Guided Attention (LGA) is proposed for efficient fusion in multi-level latent domains. Simultaneously, to effectively utilize a prior in both the k-space and image domain, under-sampled images are fused with generated full-sampled images by the Dual-domain Fusion Branch (DFB) for self-adaption guidance. Lastly, to further enhance the data consistency, we propose a k-space regularization strategy based on the non-auto-calibration signal (NACS) set. Extensive experiments on two public MRI datasets fully demonstrate the effectiveness of the proposed methodology. The code is available at <https://github.com/Zolento/MDPG>.

Keywords: Diffusion models · Generative models · MRI reconstruction.

1 Introduction

Magnetic Resonance Imaging (MRI) plays a crucial role in medical diagnostics owing to its non-invasive nature. However, its relatively time-consuming imaging process can lead to increased discomfort for the patients and pose a risk of diminished image quality. Compressed Sensing MRI (CS-MRI) reconstruction aims to recover original image signals from the under-sampled k-space acquisition, accelerating the imaging process while maintaining acceptable quality.

In recent years, generative models have been proven to outperform traditional numerical optimization algorithms in CS-MRI tasks. Recent generative adversarial networks (GANs) have been widely used in MRI reconstruction [10, 21].

Despite the great success of GANs, unstable adversarial training resulting in gradient vanishing and mode collapse is often criticized. Thus, GAN may not be able to reconstruct accurate anatomical structures [1].

Diffusion models (DMs) use a parameterized Markov chain to replace the unstable adversarial training and have been widely applied to various image inverse problems with state-of-the-art (SOTA) performance [5, 6, 18]. One step further, latent diffusion models (LDMs) [12] achieve the balance between generative quality and efficiency by transferring the diffusion process in the latent domain under a pre-trained variational autoencoder (VAE). With this feature, LDMs have become one of the most practical schemes for image generation.

Previous work has demonstrated the great potential of DMs in MRI reconstruction tasks [2, 4, 20, 23]. However, due to the random nature of DMs and the complicated anatomical structures in medical images, data consistency of DMs becomes an important factor in determining the reconstruction quality [3]. Though LDM as a novel improvement scheme can increase the efficiency of the original DM, naive data consistency in the k-space of generated images could add global artifacts. Moreover, operators in LDMs are highly non-convex, making the data consistency in the latent domain even more complicated [14, 16].

Related work has shown that the LDM’s compact prior knowledge of high dimensions helps networks learn the distribution of the original data more effectively. For instance, Li et al. [7] utilizes a prior extraction module to compress target images into a highly compact latent domain where a conditioned denoising process is applied during inference; Luo et al. [9] applies latent reconstruction error guidance in the latent domain through single-step reconstruction to improve the efficiency of deep-generated image detection; Xie et al. [19] leverages an LDM learned from high-quality face images to provide features containing prior information for blind face restoration.

Inspired by the above works, we attempt to introduce prior guidance in MRI reconstruction tasks to improve reconstruction quality. Considering the compact nature of the latent domain and the proximity of the image domain to the ground-truth distribution, we focus on explicitly aligning the generated prior guidance with measurements using auxiliary modules. Our main contributions include:

- We propose Multi-domain Diffusion Prior Guidance (MDPG), a two-stage approach integrated with prior knowledge in both latent and image domains provided by pre-trained LDMs. A novel Visual-Mamba-based backbone is utilized to process the measurements, i.e., the k-space of under-sampled images.
- We design the novel Latent Guided Attention (LGA) and Dual-domain Fusion Branch (DFB) modules to fully use prior knowledge guidance from different domains.
- We introduce a k-space regularization based on the non-auto-calibration signal (NACS) set representing the main detailed parts of images to further enhance the data consistency.

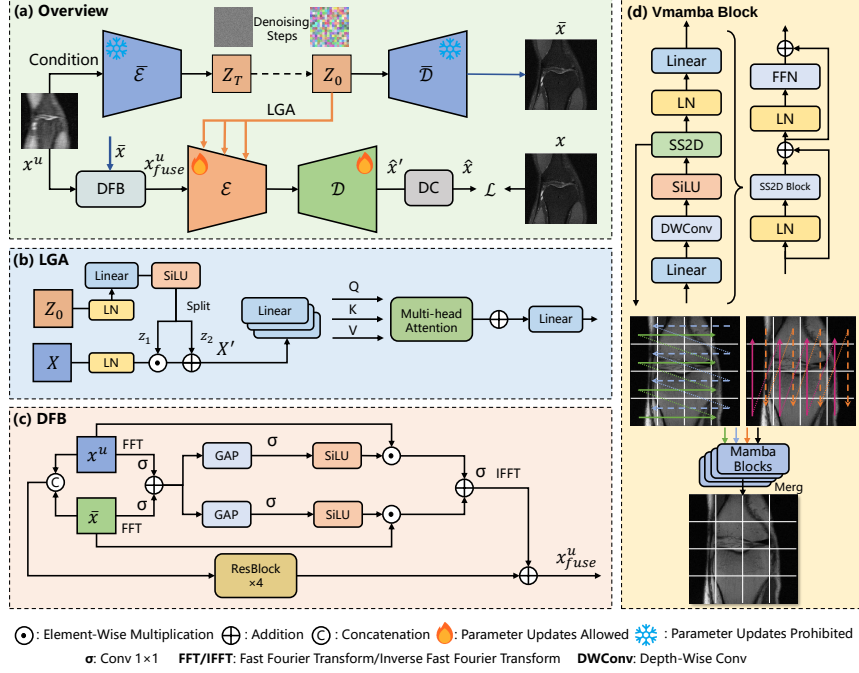


Fig. 1. Overview of the proposed methodology. (a): The overall architecture of MDPG, where $\mathcal{E}$ and $\mathcal{D}$ represent the VMamba [8] encoder and convolutional decoder of the UNet-like backbone. DC refers to the data consistency layer. (b): The basic workflow diagram for LGA. (c): The structure of DFB, where the under-sampled image x^u and synthesized full-sampled image $\hat{x}$ are fused and input into the backbone. (d): The structure and principles of VMamba blocks used in $\mathcal{E}$. SS2D refers to the 2D-Selective-Scan mechanism [8].

2 Method

We first consider the following forward sampling model for CS-MRI:

$$y = MFSx + n, n \sim N(0, \sigma_y^2 I),$$

where $x \in \mathbb{C}^N$ is the original full-sampled image, $y \in \mathbb{C}^M$ is the k-sapce measurement, and n is the measurement noise. Here, M is the under-sampling mask, $\mathcal{F}$ presents the Fourier Transform, and $\mathcal{S}$ denotes the coil sensitivity map. CS-MRI aims to reconstruct an estimated high-quality $\hat{x}$ from y :

$$\hat{x} = \arg \min_x \underbrace{\|y - MFSx\|}_{\text{data consistency}} + \underbrace{\lambda \cdot \mathcal{R}(x)}_{\text{regularization}}. \quad (1)$$

2.1 Two-stage Strategy

As illustrated in Fig.1, MDPG contains a Visual-Mamba-based backbone and a pre-trained LDM. We first train an LDM conditioned on the under-sampled image $x^u = \mathcal{F}^{-1}y$, then the backbone is trained under the guidance of LDM.

Stage I Given the under-sampled image x^u , the LDM encoder $\bar{\mathcal{E}}$ outputs a compact latent condition. We have the following conditional probability distribution [12]:

$$q(Z_T | Z) = \mathcal{N}(Z_T; \sqrt{\bar{\alpha}_T}Z, (1 - \bar{\alpha}_T)I), \quad (2)$$

where Z_t is the latent representation of t -th step, T is the total number of diffusion steps and $\bar{\alpha}_t = \prod_{i=1}^t \alpha_i$, $\alpha_t = 1 - \beta_t$ are the hyper-parameters that control the noise variance. The denoising process can be expressed as:

$$q(Z_{t-1} | Z_t, Z_0) = \mathcal{N}\left(Z_{t-1}; \mu_t(Z_t, Z_0), \frac{1 - \bar{\alpha}_{t-1}}{1 - \bar{\alpha}_t} \beta_t I\right), \quad (3)$$

$$\mu_t(Z_t, Z_0) = \frac{1}{\sqrt{\bar{\alpha}_t}}(Z_t - \frac{\beta_t}{\sqrt{1 - \bar{\alpha}_t}}\epsilon). \quad (4)$$

Finally, the predicted noise $\epsilon_\theta(Z_t, x^u, t)$ provided by a time-conditional UNet [13] is used to predict Z_{t-1} to Z_0 following the denoising process below:

$$Z_{t-1} = \frac{1}{\sqrt{\bar{\alpha}_t}}(Z_t - \frac{1 - \alpha_t}{\sqrt{1 - \bar{\alpha}_t}}\epsilon_\theta(Z_t, x^u, t)) + \sqrt{1 - \bar{\alpha}_t}\epsilon, \epsilon \sim \mathcal{N}(0, I). \quad (5)$$

We follow the original DM to train the noise-predicting UNet $\epsilon_\theta(\circ, t)$, which optimizes the following target:

$$\mathcal{L}_{\text{LDM}} := \mathbb{E}_{\bar{\mathcal{E}}(x), \epsilon \sim \mathcal{N}(0, 1), t} [\|\epsilon - \epsilon_\theta(Z_t, x^u, t)\|_2^2]. \quad (6)$$

Stage II Once the LDM is trained, DDIM sampler [17] and LDM decoder $\bar{\mathcal{D}}$ are used to generate Z_0 and synthesized full-sampled image $\bar{x}$. Given that down-sampling can create global artifacts in the image domain, a larger receptive field aids in feature extraction and fusion. Leveraging the innovative selective scan mechanism, the recently introduced VMamba model can extract global features from images with linear complexity. Consequently, we utilize VMamba for efficient processing of Z_0 and $\bar{x}$ within the encoder. Conversely, we opt for a straightforward convolutional design to restore fine-grained details while minimizing computational overhead, thereby providing sufficient precision in the decoder.

As Fig.1 demonstrates, Z_0 and $\bar{x}$ are integrated with the VMamba encoder $\mathcal{E}$ by the LGA and DFB modules, respectively. The architecture of the backbone follows the UNet [13], where the convolution layers in the original encoder are replaced with VMamba layers which follow the design in Fig.1(d). The final loss function is described in Eq.13.

2.2 Prior Knowledge Guidance

The compact prior Z_0 and corresponding decoded synthesized full-sampled image $\bar{x}$ are given during the training process of the backbone. For Z_0 , LGA is proposed to inject prior knowledge into the backbone encoder $\mathcal{E}$. Let W denote linear projection layers, LN refer to the layer-norm operation, and X represent the feature in the corresponding layer of $\mathcal{E}$, LGA can be expressed as:

$$\begin{aligned} \text{LGA}(Q, K, V) &= \text{Softmax}\left(\frac{QK^T}{\sqrt{d_k}}\right)V, \\ \text{where } Q, K, V &= X'W^Q, X'W^K, X'W^V, X' = \text{LN}(X) \cdot z_1 + z_2, \\ z_1, z_2 &= \text{Split}(\text{SiLU}(\text{LN}(Z_0)W^Z)). \end{aligned} \quad (7)$$

The k-space as the measurement space of MRI signals contains a wealth of original frequency information, which helps to enhance data consistency. Thus, we design DFB to extract the rich details from both the image domain and frequency domain, as shown in Fig.1. Considering the input image's under-sampled characteristic in the k-space, a dynamic weighting strategy is chosen for complementarity. Meanwhile, we use a residual branch to extract information from the concatenated image features. The detailed process is as follows:

$$x_{fuse}^u = \sigma_{iout}(\mathcal{F}^{-1}[f(\sigma_x(x_k^u), \sigma_y(\bar{x}_k)) + g(x^u, \bar{x})]), \quad (8)$$

$$f(x, y) = \sigma_{kout}[x \cdot \text{SiLU}(t_1(x, y)) + y \cdot \text{SiLU}(t_2(x, y))], \quad (9)$$

$$t_i(x, y) = \sigma_i(\phi(x + y)), \quad (10)$$

$$g(x, y) = \text{ResBlocks}(\text{concat}(x, y)), \quad (11)$$

where $\sigma(\cdot)$ represents the 1×1 convolution and $\phi(\cdot)$ denotes the global average pooling (GAP) operation, respectively. The subscript k indicates that the feature is transformed to the frequency domain, i.e., k-space. The output features fused by the DFB are fed again into the encoder $\mathcal{E}$ as input to the backbone network.

2.3 NACS Set based K-space Regularization

Degradation modes in under-sampled images are closely related to the under-sampling mask. The data consistency layer is often achieved by a learnable parameter ν related to noise level n [15]:

$$\hat{x}_k(i_k, j_k) = \begin{cases} \hat{x}_k'(i_k, j_k), & (i_k, j_k) \notin \Omega \\ \frac{\hat{x}_k'(i_k, j_k) + \nu x_k^u(i_k, j_k)}{1 + \nu}, & (i_k, j_k) \in \Omega \end{cases} \quad (12)$$

Ω represents the under-sampled k-space set and $\hat{x}_k'$ denotes k-space of the original output of the backbone. This approach results in the optimization being dominated by the main low-frequency signal represented by the ACS. We recommend optimizing k-space separately according to the NACS set which represents most

Table 1. Reconstruction results of FastMRI dataset. Best results are bolded while second best are underlined.

Method	8× 0.04c			4× 0.08c		
	PSNR↑	SSIM↑	NMSE↓	PSNR↑	SSIM↑	NMSE↓
ZF (Zero-Filling)	26.82	0.619	0.0893	29.47	0.729	0.0531
ISTA-Net [22]	<u>29.75</u>	0.724	<u>0.0502</u>	31.54	0.759	0.0410
DCCNN [15]	29.25	0.710	0.0546	<u>31.97</u>	0.787	0.0346
MD-Recon [15]	29.52	0.717	0.0521	31.94	<u>0.789</u>	<u>0.0343</u>
UNet [13]	29.51	<u>0.725</u>	0.0546	31.68	0.787	0.0382
DAGAN [21]	28.40	0.684	0.0635	30.91	0.766	0.0407
DDNM [18]	27.89	0.665	0.0833	31.35	0.762	0.0418
LDM-16 [12]	25.99	0.536	0.113	26.98	0.578	0.0945
MC-DDPM [20]	27.35	0.614	0.0900	29.16	0.683	0.0658
MDPG (ours)	30.43	0.740	0.0451	32.14	0.794	0.0334

Table 2. Reconstruction results of IXI T1 dataset. The middle 10 slices of each volume were used to calculate the metrics. Best results are bolded while second best are underlined.

Method	8× 0.04c			4× 0.08c		
	PSNR↑	SSIM↑	NMSE↓	PSNR↑	SSIM↑	NMSE↓
ZF (Zero-Filling)	21.75	0.529	0.103	24.45	0.668	0.0554
ISTA-Net [22]	<u>26.97</u>	<u>0.829</u>	<u>0.0313</u>	32.37	0.942	0.00915
MD-Recon [11]	26.90	0.813	0.0317	<u>32.01</u>	<u>0.932</u>	<u>0.00988</u>
UNet [13]	26.47	0.796	0.0353	30.06	0.891	0.0155
DAGAN [21]	24.30	0.735	0.0573	25.86	0.787	0.0400
DDNM [18]	24.11	0.747	0.0607	29.10	0.870	0.0191
LDM-16 [12]	20.99	0.600	0.122	22.56	0.679	0.0854
MC-DDPM [20]	23.97	0.672	0.0624	28.21	0.800	0.0239
MDPG (ours)	27.50	0.833	0.0278	31.00	0.914	0.0124

of the high-frequency, i.e., detailed parts of images. The combined loss function is defined as:

$$\mathcal{L} = \frac{\|\hat{x} - x\|_2}{\|x\|_2} + \lambda_1 \frac{\|\hat{x}_k^{\Gamma} - x_k^{\Gamma}\|_2}{\|x_k^{\Gamma}\|_2} \quad (13)$$

The first term denotes minimizing the loss in the image domain. The second term regularizes the k-space according to the sampling pattern. Here Γ denotes the NACS set. λ_1 is a fixed hyperparameter.

3 Experiments and Results

3.1 Datasets and Baselines

We evaluate our method on two MRI datasets: FastMRI¹ single-coil knee dataset and IXI² T1 dataset. The FastMRI dataset consists of a training set with 973

¹ <https://fastmri.med.nyu.edu/>

² <https://brain-development.org/ixi-dataset/>

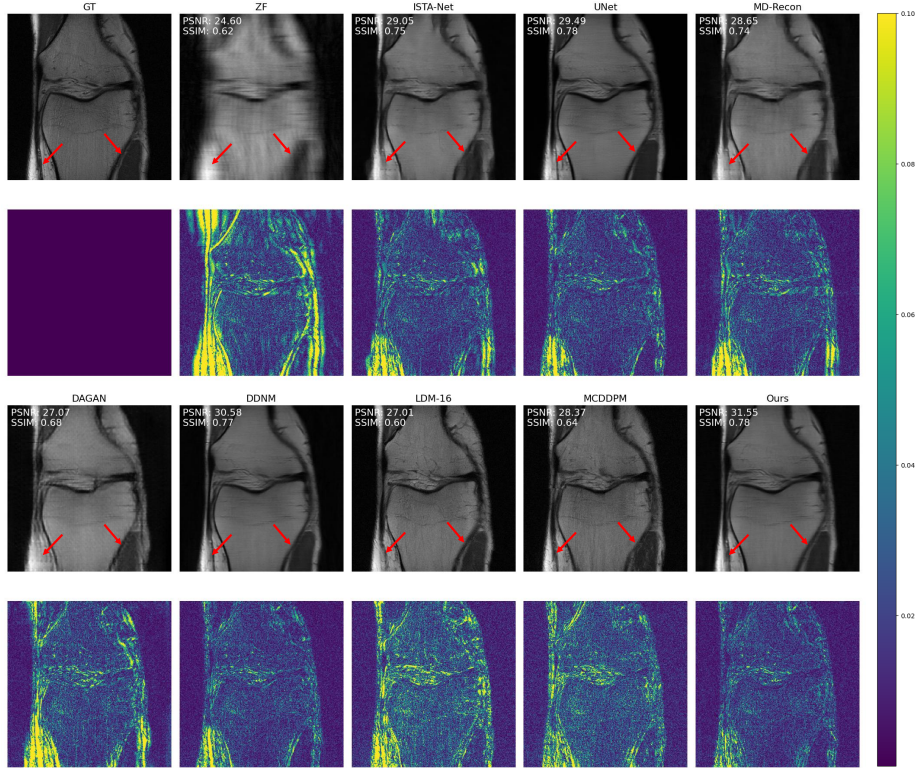


Fig. 2. Visualization of reconstruction results and corresponding error maps of FastMRI dataset of $8\times$ acceleration factor. Red arrows indicate highly degraded key anatomical structures that are desired to be as close as possible to the ground truth (GT).

volumes and a validation set with 199 volumes. The IXI T1 dataset is divided into a training set with 462 volumes and a validation set with 115 volumes. For preprocessing, each of the slices in the FastMRI dataset is cropped into 320×320 size, while those in the IXI T1 dataset are cropped into 256×256 . Under-sampled k-space was generated by applying $4\times$ and $8\times$ Cartesian masks on full-sampled k-space using official codes of FastMRI.

We compare our method with different types of MRI reconstruction baselines. For end-to-end methods, we choose ISTA-Net [22], UNet [13], and MD-Recon-Net [11]. For generative models, we choose DAGAN [21], LDM [12], DDNM [18], and MC-DDPM [20]. Average metrics are computed on every volume of the validation set, including Peak Signal-to-Noise Ratio (PSNR), Structural Similarity Index Measure (SSIM), and Normalized Mean Square Error (NMSE) values.

Table 3. Ablation results of components and training configurations in MDPG under $8\times$ acceleration of the FastMRI dataset. A: DFB; B: LGA; C: using Sigmoid instead of SiLU in DFB; D: learnable ν ; E: placing DFB before $\mathcal{E}$ instead of after $\mathcal{D}$. F: NACS set regularization. The first row represents the results of the backbone without prior guidance.

A	B	C	D	E	F	PSNR $\uparrow$	SSIM $\uparrow$	NMSE $\downarrow$
N/A	N/A	N/A	✓	N/A	✗	29.97	0.731	0.04870
✓	✓	✗	✓	✗	✗	29.56	0.679	0.05776
✓	✓	✓	✓	✗	✗	30.10	0.730	0.04812
✓	✓	✓	✗	✓	✗	30.23	0.735	0.04700
✓	✓	✗	✓	✓	✗	30.12	0.733	0.04745
✓	✗	✓	✓	✓	✗	30.20	0.736	0.04676
✗	✓	✓	✓	✓	✗	30.24	0.738	0.04627
✓	✓	✓	✓	✓	✗	<u>30.33</u>	<u>0.739</u>	<u>0.04614</u>
✓	✓	✓	✓	✓	✓	30.43	0.740	0.04514

3.2 Implementation Details

We trained MDPG’s backbone with learning rate = 4×10^{-3} , batch size = 8, and an AdamW optimizer for 100 epochs. LDM was trained with $T = 1000$, $16\times$ down-sampled Z , and a linear noise schedule. The sampling step of DDIM is set to 20 in both the training and inference processes. λ_1 is set to 0.1. All experiments are done with an NVIDIA GeForce RTX 4090 GPU under the Pytorch framework.

3.3 Results

As Table 1 and 2 demonstrate, our method achieves the best performance in three scenarios. In the complex-valued FastMRI dataset, MDPG has universal advantages in all metrics. Fig.2 demonstrates how our method outperforms other baselines. MDPG restores two highly degraded anatomical structures more accurately than end-to-end baselines without introducing excessive false details like generative models. In other regions, MDPG also demonstrates excellent performance without exhibiting significant errors. For the real-valued IXI T1 dataset, MDPG is more advantageous under $8\times$ acceleration. Though it does not yield optimal results at $4\times$ factors, the performance was still better than all generative baselines. The performance degradation is most likely attributed to the low quality of the generated prior, as evidenced by the low metrics for LDM under $4\times$ acceleration.

3.4 Ablation Study

We conducted ablation experiments to verify the effectiveness of each key component and configuration in MDPG. As Table 3 shows, priors in the latent, image, and frequency domains improve performance over the backbone. Experiments also demonstrate performance improvement by the NACS set regularization.

4 Conclusion

This paper introduces MDPG, a novel approach incorporating a latent diffusion prior, a multi-domain fusion strategy, and k-space regularization based on the NACS set. The experiments demonstrate that our method outperforms other baseline approaches. In conclusion, MDPG illustrates that the utilization of multi-domain diffusion prior knowledge could effectively enhance the performance of MRI reconstruction tasks.

Disclosure of Interests. The authors have no competing interests to declare that are relevant to the content of this article.

References

1. Bau, D., Zhu, J.Y., Wulff, J., Peebles, W., Strobel, H., Zhou, B., Torralba, A.: Seeing what a gan cannot generate. In: Proceedings of the IEEE/CVF international conference on computer vision. pp. 4502–4511 (2019)
2. Bian, W., Jang, A., Zhang, L., Yang, X., Stewart, Z., Liu, F.: Diffusion modeling with domain-conditioned prior guidance for accelerated mri and qmri reconstruction. *IEEE Transactions on Medical Imaging* (2024)
3. Chung, H., Kim, J., Mccann, M.T., Klasky, M.L., Ye, J.C.: Diffusion posterior sampling for general noisy inverse problems. In: The Eleventh International Conference on Learning Representations (2023), <https://openreview.net/forum?id=OnD9zGAGT0k>
4. Gao, Z., Zhou, S.K.: U²MRPD: Unsupervised undersampled mri reconstruction by prompting a large latent diffusion model. *arXiv preprint arXiv:2402.10609* (2024)
5. Ho, J., Jain, A., Abbeel, P.: Denoising diffusion probabilistic models. *Advances in neural information processing systems* **33**, 6840–6851 (2020)
6. Kavar, B., Elad, M., Ermon, S., Song, J.: Denoising diffusion restoration models. *Advances in Neural Information Processing Systems* **35**, 23593–23606 (2022)
7. Li, G., Rao, C., Mo, J., Zhang, Z., Xing, W., Zhao, L.: Rethinking diffusion model for multi-contrast mri super-resolution. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 11365–11374 (2024)
8. Liu, Y., Tian, Y., Zhao, Y., Yu, H., Xie, L., Wang, Y., Ye, Q., Jiao, J., Liu, Y.: Vmamba: Visual state space model. *Advances in neural information processing systems* **37**, 103031–103063 (2024)
9. Luo, Y., Du, J., Yan, K., Ding, S.: Lare²: Latent reconstruction error based method for diffusion-generated image detection. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 17006–17015 (2024)
10. Noor, R., Wahid, A., Bazai, S.U., Khan, A., Fang, M., Syam, M., Bhatti, U.A., Ghadi, Y.Y.: Dlgan: Undersampled mri reconstruction using deep learning based generative adversarial network. *Biomedical Signal Processing and Control* **93**, 106218 (2024)
11. Ran, M., Xia, W., Huang, Y., Lu, Z., Bao, P., Liu, Y., Sun, H., Zhou, J., Zhang, Y.: Md-recon-net: a parallel dual-domain convolutional neural network for compressed sensing mri. *IEEE Transactions on Radiation and Plasma Medical Sciences* **5**(1), 120–135 (2020)
12. Rombach, R., Blattmann, A., Lorenz, D., Esser, P., Ommer, B.: High-resolution image synthesis with latent diffusion models. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 10684–10695 (2022)

13. Ronneberger, O., Fischer, P., Brox, T.: U-net: Convolutional networks for biomedical image segmentation. In: Medical image computing and computer-assisted intervention—MICCAI 2015: 18th international conference, Munich, Germany, October 5–9, 2015, proceedings, part III 18. pp. 234–241. Springer (2015)
14. Rout, L., Raoof, N., Daras, G., Caramanis, C., Dimakis, A., Shakkottai, S.: Solving linear inverse problems provably via posterior sampling with latent diffusion models. *Advances in Neural Information Processing Systems* **36** (2024)
15. Schlemper, J., Caballero, J., Hajnal, J.V., Price, A.N., Rueckert, D.: A deep cascade of convolutional neural networks for dynamic mr image reconstruction. *IEEE transactions on Medical Imaging* **37**(2), 491–503 (2017)
16. Song, B., Kwon, S.M., Zhang, Z., Hu, X., Qu, Q., Shen, L.: Solving inverse problems with latent diffusion models via hard data consistency. In: The Twelfth International Conference on Learning Representations (2024), <https://openreview.net/forum?id=j8hdRqOUhN>
17. Song, J., Meng, C., Ermon, S.: Denoising diffusion implicit models. In: International Conference on Learning Representations (2021), <https://openreview.net/forum?id=St1giarCHLP>
18. Wang, Y., Yu, J., Zhang, J.: Zero-shot image restoration using denoising diffusion null-space model. In: The Eleventh International Conference on Learning Representations (2023), <https://openreview.net/forum?id=mRieQgMtNTQ>
19. Xie, L., Zheng, C., Xue, W., Jiang, L., Liu, C., Wu, S., Wong, H.S.: Learning degradation-unaware representation with prior-based latent transformations for blind face restoration. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 9120–9129 (2024)
20. Xie, Y., Li, Q.: Measurement-conditioned denoising diffusion probabilistic model for under-sampled medical image reconstruction. In: International Conference on Medical Image Computing and Computer-Assisted Intervention. pp. 655–664. Springer (2022)
21. Yang, G., Yu, S., Dong, H., Slabaugh, G., Dragotti, P.L., Ye, X., Liu, F., Arridge, S., Keegan, J., Guo, Y., et al.: Dagan: deep de-aliasing generative adversarial networks for fast compressed sensing mri reconstruction. *IEEE transactions on medical imaging* **37**(6), 1310–1321 (2017)
22. Zhang, J., Ghanem, B.: Ista-net: Interpretable optimization-inspired deep network for image compressive sensing. In: Proceedings of the IEEE conference on computer vision and pattern recognition. pp. 1828–1837 (2018)
23. Zhao, X., Yang, T., Li, B., Yang, A., Yan, Y., Jiao, C.: Diffgan: An adversarial diffusion model with local transformer for mri reconstruction. *Magnetic Resonance Imaging* **109**, 108–119 (2024)