

MR-DiffRecon: Efficient All-in-One MRI Reconstruction with Multi-Scale Alignment and Implicit Prompting

Wafa Al Ghallabi^{1*}, Akshay Dudhane¹, Salman Khan¹, and Fahad Shahbaz Khan^{1,2}

¹ Mohamed bin Zayed University of Artificial Intelligence, Abu Dhabi, UAE

² Linköping University, Linköping, Sweden

Abstract. Magnetic Resonance Imaging (MRI) offers unmatched soft-tissue contrast without ionizing radiation, yet is among the least equitably distributed modalities: long acquisitions limit patient throughput, while modern reconstruction networks can exceed the compute available in low-resource clinics. Accelerated reconstruction can reduce acquisition burden, but many high-performing unrolled and diffusion models are large and trained separately for different acquisition settings. We revisit accelerated MRI reconstruction through the lens of resource efficiency and propose **MR-DiffRecon**, a unified, single-stage framework that supports highly undersampled acquisitions while reducing the inference footprint. Within each dataset, a single model reconstructs across contrasts, views, and multiple acceleration factors (up to $\times 10$), while the same architectural design is independently validated on a second anatomy using dataset-specific training. Adjacent-slice features are aligned at multiple scales via deformable convolutions and fused with implicitly learned prompts, while a training-time reverse diffusion process improves fidelity without requiring diffusion sampling at inference. MR-DiffRecon achieves the strongest SSIM and PSNR among the evaluated methods on CMRxRecon and FastMRI multi-coil knee, while using 23.49% fewer parameters and 7.7% fewer GFLOPs than PromptMR. Code and models are available at <https://github.com/wafaAlghallabi/MR-DiffRecon>.

Keywords: Resource-constrained settings · Accelerated MRI reconstruction · Efficient deep learning · Denoising diffusion probabilistic model.

1 Introduction

Magnetic resonance imaging (MRI) is a cornerstone of modern diagnostics, offering detailed, radiation-free insight into anatomy and pathology, yet access is profoundly unequal. Many low- and middle-income countries operate only a handful of scanners for millions of people, and where hardware exists it is often older, lower-field, and supported by limited computing and power infrastructure [5, 13]. Two coupled bottlenecks drive this scarcity: acquisition is slow,

* Corresponding author: wafa.alghallabi@mbzuai.ac.ae

limiting scanner throughput, while state-of-the-art reconstruction networks can be computationally demanding [5]. A method intended for real-world impact must therefore be judged not only on fidelity but also on its acquisition and inference requirements.

Accelerated reconstruction addresses the acquisition bottleneck by recovering images from undersampled k-space, and deep learning now dominates this task (Section 2). These gains, however, often come with large models trained separately for different anatomies, contrasts, or rates, increasing storage and maintenance burden. Accelerated reconstruction should therefore be evaluated not only for fidelity, but also for the accuracy–efficiency trade-off relevant to deployment under constraint.

To this end we propose **MR-DiffRecon**, a single-stage inference, all-in-one framework coupling unrolled CNN feature extraction with training-time diffusion refinement. It aligns adjacent slices via multi-scale deformable convolutions, fuses them with input-adaptive implicit prompts, and uses reverse diffusion during training to improve fidelity without an explicit refinement network at inference. Our contributions are threefold:

- **Access-oriented reconstruction.** We present an efficient all-in-one model that supports highly accelerated MRI reconstruction while reducing the inference footprint relative to comparable all-in-one baselines.
- **Unified within-dataset design.** Multi-scale feature alignment corrects inter-slice misalignment, implicit prompting lets one model span views, contrasts, and undersampling rates within each dataset, and the same architecture is validated independently on cardiac and knee MRI using dataset-specific training.
- **Accuracy at a low footprint.** MR-DiffRecon achieves the strongest SSIM and PSNR among the evaluated methods on CMRxRecon and FastMRI knee with 23.49% fewer parameters and 7.7% fewer GFLOPs than PromptMR; its training-time diffusion refinement adds no diffusion sampling cost at deployment.

2 Previous Work

Classical and unrolled reconstruction. Compressed-sensing methods [11] recover images from undersampled k-space by exploiting sparsity but trade fidelity against speed and need per-scan tuning [2]. Deep learning has since taken over via unrolled networks that alternate learned regularization with data consistency [8]. E2E-VarNet [15] embeds coil-sensitivity estimation and data consistency end-to-end and remains a strong baseline, while HUMUS-Net [4] raises fidelity with CNN–transformer hybrids at a steep parameter and compute cost. Both are typically trained separately per anatomy, contrast, and rate.

Diffusion-based reconstruction. Denoising diffusion models deliver strong fidelity for MRI reconstruction and related restoration [6, 10, 12, 18], but iterative

sampling increases inference latency and compute. MR-DiffRecon uses diffusion only during training and therefore avoids iterative diffusion sampling at deployment (Section 3).

All-in-one and prompt-based restoration. PromptIR [14] conditions an image-restoration network on input-specific prompts, enabling all-in-one operation; PromptMR [19] extends this idea to dynamic and multi-contrast MRI using adjacent slices and is our most directly comparable baseline. PromptMR relies on an external refinement network and does not explicitly estimate and correct inter-slice displacement. Across these families, progress is measured primarily by fidelity, while model size, compute, and the burden of maintaining many task-specific networks receive less attention [7, 17, 20, 22]. MR-DiffRecon targets this gap by combining alignment, prompting, and training-time diffusion refinement with a compact single-stage inference model.

3 Methodology

MR-DiffRecon integrates iterative unrolling, multi-scale feature alignment, adaptive prompting, and training-time reverse diffusion in a single stage, handling diverse undersampling patterns ($\times 4$, $\times 8$, $\times 10$) and contrasts at reduced inference cost. Unlike two-stage models such as PromptMR [19], it embeds fusion and prompting *inside* the iterative structure and removes diffusion refinement at inference.

3.1 Iterative Unrolling Backbone

We adopt the standard unrolled formulation for recovering a complex image x from undersampled multi-coil k-space $y = Ax + \epsilon$, where $A = MFS$ couples coil-sensitivity encoding S , Fourier transform F , and undersampling mask M . Reconstruction minimizes $\frac{1}{2}\|Ax - y\|_2^2 + \lambda R(x)$, unrolled into a cascade of gradient-style updates with learned step sizes and a learned image-domain regularizer, following established practice [8, 15]. Sensitivity estimation and data consistency are incorporated into the unrolled reconstruction, and the final image is formed by inverse Fourier transform followed by root-sum-of-squares (RSS) aggregation. We use 8 cascades throughout to reduce inference complexity relative to the 12-cascade PromptMR configuration [19]; the component ablation in Table 4 is evaluated under this fixed backbone. As shown in Fig. 1, each cascade is a four-level encoder-decoder exploiting spatial and anatomical redundancy.

Dual Attention Block (DAB). The DAB combines spatial attention (local saliency) and channel attention (global reweighting), applied symmetrically across encoder and decoder; encoder features are downsampled by strided convolution, decoder features upsampled by transposed convolution. It sharpens features without increasing the cascade width.

Multi-Scale Feature Alignment. Adjacent slices may be misaligned by patient motion or acquisition inconsistency. We correct this with deformable convolution [1], following the multi-frame alignment strategy of [3]: offsets predicted

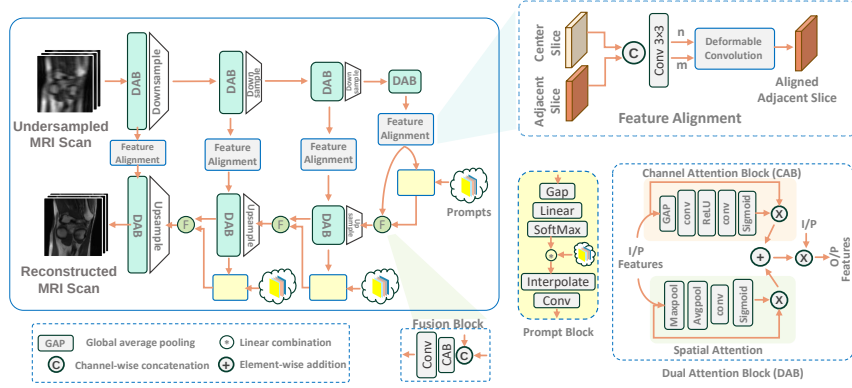


Fig. 1: **MR-DiffRecon architecture.** Each cascade applies deformable multi-scale alignment, dual-attention refinement, and implicit prompting. A lightweight fusion block merges aligned adjacent-slice features with center-slice features. Reverse diffusion is used only during training; inference uses the 8-cascade backbone alone.

from the center- and adjacent-slice features adapt the sampling grid so that adjacent-slice features register to the center slice. For center- and adjacent-slice features $\mathbf{F}^t, \mathbf{F}^i$,

$$\Delta m, \Delta n = W^o(\mathbf{F}^t, \mathbf{F}^i), \quad \mathbf{A}_F = W^d(\mathbf{F}^i, \Delta m, \Delta n), \quad (1)$$

with learnable offset/deformable convolutions W^o, W^d , producing the aligned adjacent-slice feature $\mathbf{A}_F$. Applying this at each of the four encoder-decoder levels lets the model exploit adjacent-slice information more reliably.

Fusion Block. At each scale, the aligned adjacent-slice feature $\mathbf{A}_F$ is merged with the center-slice feature by channel-wise concatenation followed by a 1×1 convolution and a channel-attention block (CAB), which reweights the concatenated channels so that reliably aligned adjacent-slice information is emphasized and residual misalignment is suppressed. The fused feature is passed to the decoder.

Implicit Prompt Block. Following PromptIR [14], an implicit prompt conditions the decoder on input-specific context, enabling a single model to serve multiple acquisition settings within a dataset. For decoder feature $F_{d,i}$,

$$v = \text{Linear}(\text{GAP}(F_{d,i})), \quad p = \text{Softmax}(v), \quad I_p = \text{Interp}(p), \quad \hat{P}_i = W_p \odot F_{d,i}, \quad (2)$$

where W_p is a learned projection of the interpolated prompt I_p .

3.2 Reverse Diffusion Refinement

To improve reconstruction fidelity during training, we apply a reverse diffusion process to the final 2D reconstruction produced by the unrolled backbone. After the eighth cascade, the multi-coil output is transformed to the image domain by inverse Fourier transform (iFFT) and combined using root-sum-of-squares (RSS),

yielding the reconstructed slice $\mathbf{I}$. Following [12], the reverse process is modeled as

$$p_{\theta}(\mathbf{I}_{0:T}) = p(\mathbf{I}_T) \prod_{t=1}^T \mathcal{N}(\mathbf{I}_{t-1}; \mu_{\theta}(\mathbf{I}_t, t), \sigma_{\theta}^2(\mathbf{I}_t, t)), \quad (3)$$

where $\mathbf{I}_T$ denotes the noisy state and $\mathbf{I}_0$ the refined reconstruction. The terms $\mu_{\theta}(\mathbf{I}_t, t)$ and $\sigma_{\theta}^2(\mathbf{I}_t, t)$ denote the learned mean and variance of the reverse transition at step t . Successive reverse steps progressively reduce noise in the reconstructed slice. This refinement is used only during training; at inference it is removed entirely and reconstruction is performed by the 8-cascade unrolled backbone alone, introducing no diffusion-related inference parameters or sampling steps.

Datasets and Splits: We evaluate on two public multi-coil datasets, following the data preparation and splits of PromptMR [19]. CMRxRecon [16] provides 120 cardiac cases (3T) with dynamic cine (short-axis, SAX; 2-/3-/4-chamber long-axis, LAX) and multi-contrast T1w/T2w mapping, with k-space compressed to 10 virtual coils. We split the 120 cases into 100 training and 20 validation cases and report on the validation set, as no public test labels are available. Cardiac data use an equispaced mask with 24 fixed low-frequency lines at acceleration factors $\times 4$, $\times 8$, and $\times 10$; results are reported at $\times 10$. For FastMRI knee [21] we train on the official multi-coil knee training set (973 volumes) and, because the online test server is retired, split the 199 official validation volumes into 99 validation and 100 test volumes following [19]; results are reported on the 100-case test set. Knee data use a random mask at acceleration factors $\times 4$ and $\times 8$; results are reported at $\times 8$. Each input contains the center slice t and four neighbors ($t \pm 1$, $t \pm 2$).

Implementation and Compute Resources: All models are trained with Adam [9] (learning rate 10^{-4} , weight decay 0.01) using an SSIM loss; every MR-DiffRecon instance uses 8 cascades. We train separate instances per dataset for 100 epochs on CMRxRecon and 45 epochs on FastMRI with a batch size of 8, on $4 \times$ NVIDIA A100 GPUs, taking approximately two days per dataset. Training is a one-time, centralized cost, decoupled from deployment: our efficiency claim concerns inference, where only the 8-cascade backbone is used. All baselines were trained and evaluated under matched settings for a fair comparison.

Efficiency Protocol: We report inference cost as parameter count and per-slice FLOPs at the CMRxRecon 10-coil setting. SSIM, PSNR, and NMSE are computed independently under identical preprocessing and aggregation settings; therefore, relative rankings across metrics need not coincide.

4 Results

Reconstruction Accuracy: As shown in Table 1, MR-DiffRecon achieves the best SSIM and PSNR among the compared methods. Notably, our *all-in-one* model also outperforms the *one-by-one* baselines trained separately per task

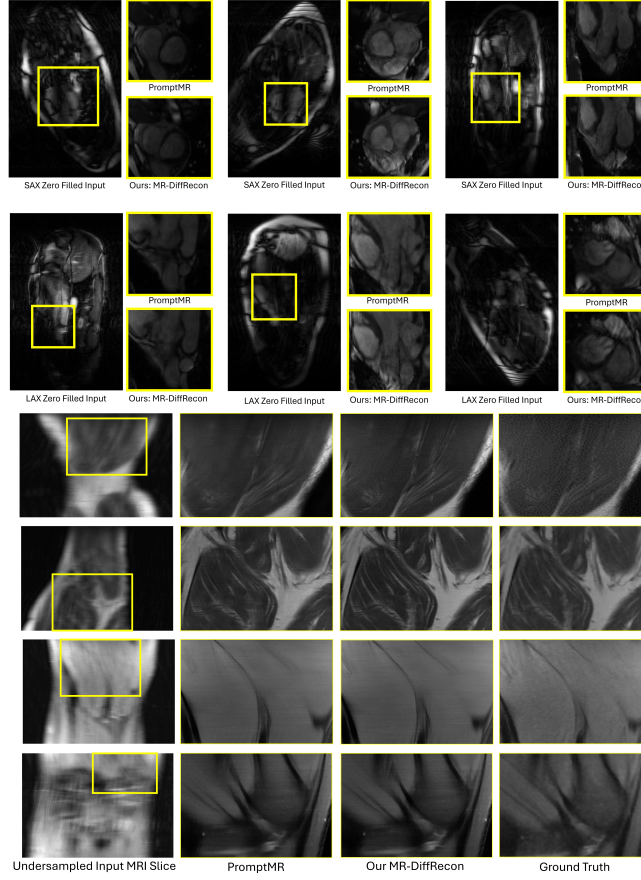


Fig. 2: **Qualitative reconstruction.** *Top:* CMRxRecon at $\times 10$ (zero-filled input, PromptMR [19], MR-DiffRecon). *Bottom:* FastMRI knee at $\times 8$ (undersampled input, PromptMR, MR-DiffRecon, ground truth).

(E2E-VarNet, HUMUS-Net-L). Relative to PromptMR, MR-DiffRecon gains +1.05 dB on Cine SAX and +1.22 dB on Cine LAX; relative to PromptMR+Shift, the corresponding improvements are +1.00 dB and +1.18 dB, while using 8 rather than 12 cascades. On FastMRI knee (Table 2), MR-DiffRecon reaches 37.95 dB PSNR with an SSIM improvement from 0.8970 to 0.9111 over PromptMR, demonstrating that the same architectural design remains effective on a second anatomy after dataset-specific training. Fig. 2 provides qualitative comparisons with PromptMR on both datasets.

Efficiency and Resource Footprint: As shown in Table 3 and Fig. 3, MR-DiffRecon has the lowest parameter count of all methods (53.33M) and the lowest FLOPs among all-in-one models, using 23.49% fewer parameters and 7.7% fewer GFLOPs than PromptMR, and far less compute than PromptIR (4455 GFLOPs). The one-by-one E2E-VarNet and HUMUS-Net-L use fewer FLOPs

Table 1: Reconstruction accuracy on CMRxRecon (validation set, $\times 10$), reported as SSIM/PSNR/NMSE($\times 10^{-2}$). MR-DiffRecon and HUMUS-Net-L use 8 cascades; other unrolled methods use 12. PromptMR+Shift is the slice-shift-augmented variant [19].

Task	Method	Cine		Mapping	
		SAX	LAX	T1w	T2w
One-by-One	E2E-VarNet [15]	0.9744/42.05/1.6	0.9673/39.93/2.1	0.98/43.12/1.5	0.9777/41.20/1.4
	HUMUS-Net-L [4]	0.9791/42.96/1.3	0.9689/40.07/2.0	0.9832/43.85/1.3	0.9824/42.39/1.1
All-in-One	PromptIR [14]	0.9659/40.16/2.5	0.9581/38.62/2.7	0.9726/41.10/2.3	0.9784/41.10/1.4
	PromptMR [19]	0.9865/45.58/1.1	0.9836/43.72/1.2	0.9899/46.84/1.0	0.9903/46.24/0.7
	PromptMR+Shift [19]	0.9866/45.63/0.7	0.9837/43.76/0.9	0.9903/47.04/0.7	0.9905/46.33/0.5
All-in-One	MR-DiffRecon	0.9947/46.63/0.7	0.9987/44.94/0.9	0.9982/47.23/0.7	0.9982/47.06/0.5

Table 2: Reconstruction accuracy on the FastMRI knee multi-coil test set (100 cases, $\times 8$).

Method	SSIM $\uparrow$	PSNR $\uparrow$	NMSE($\times 10^{-2}$) $\downarrow$
E2E-VarNet [15]	0.8936 $\pm$ 0.1157	37.30 $\pm$ 4.925	0.8690 $\pm$ 0.9279
HUMUS-Net [4]	0.8946 $\pm$ 0.1162	37.20 $\pm$ 5.009	0.8974 $\pm$ 0.9743
HUMUS-Net-L [4]	0.8955 $\pm$ 0.1161	37.45 $\pm$ 5.067	0.8587 $\pm$ 0.9930
PromptMR [19]	0.8970 $\pm$ 0.1168	37.63 $\pm$ 5.319	0.8344 $\pm$ 0.9648
MR-DiffRecon	0.9111 $\pm$ 0.1152	37.95 $\pm$ 5.193	0.8247 $\pm$ 0.9447

but reconstruct a single task at a time and trail in accuracy. Because diffusion is removed after training, the reported MR-DiffRecon FLOPs correspond to the deployed 8-cascade backbone.

Ablation: On $\times 10$ Cine SAX (Table 4), the cumulative ablation shows monotonic improvement as each component is introduced. Replacing PromptMR’s channel attention with our DAB gives 45.93 dB PSNR; adding single- then multi-scale alignment raises it to 46.06 and 46.18 dB, and adding reverse diffusion reaches 46.63 dB. Together the modules yield a 0.70 dB improvement over the DAB-only baseline, with corresponding gains in SSIM and NMSE.

5 Discussion

Benefits of within-dataset consolidation. A single all-in-one MR-DiffRecon surpasses task-specific baselines on CMRxRecon (Table 1) while jointly handling multiple views, contrasts, and acceleration factors. The cumulative ablation suggests that implicit prompting and multi-scale deformable alignment support this behavior by adapting shared features to the input and improving the use of adjacent-slice information. The same architectural design also remains effective on FastMRI knee after dataset-specific training, indicating transferability of the design rather than a single anatomy-agnostic model.

Refinement without an inference penalty. The largest incremental improvement in the cumulative ablation occurs when training-time reverse diffusion

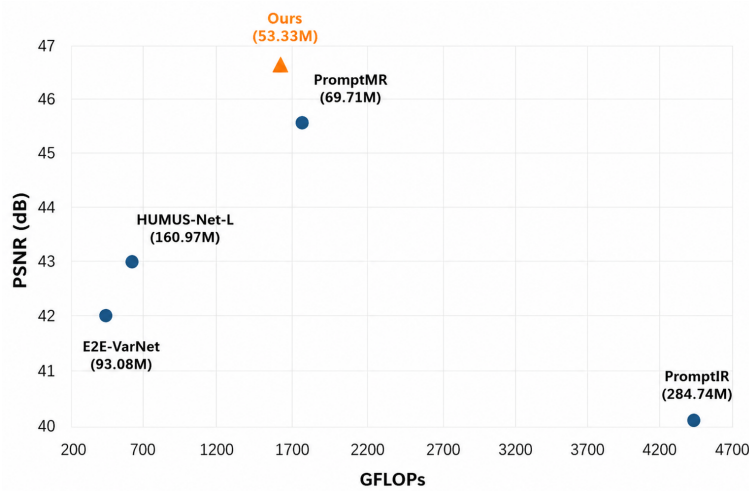


Fig. 3: **Accuracy vs. compute** on CMRxRecon Cine SAX ($\times 10$). MR-DiffRecon (orange) achieves the highest PSNR with the lowest parameter count (53.33M) among the compared methods while maintaining a relatively low inference FLOP budget. Marker labels indicate parameter counts.

is introduced, increasing Cine SAX PSNR from 46.18 to 46.63 dB and SSIM from 0.9891 to 0.9947 (Table 4). Since diffusion is removed after training, the deployed model remains the 8-cascade backbone and requires no iterative diffusion sampling. Because the backbone depth is fixed across the ablation, these improvements are measured under a common reconstruction depth.

6 Impact in Resource-Constrained Settings

Throughput and deployment. Conventional MRI examinations can be lengthy [5]; reconstruction from highly undersampled acquisitions ($\times 8$ – $\times 10$ in Tables 1, 2) has the potential to reduce the acquisition time of accelerated sequences, subject to prospective implementation and clinical validation. MR-DiffRecon also reduces parameter count and all-in-one inference FLOPs relative to PromptMR (Table 3), lowering the deployment burden compared with maintaining multiple acquisition-specific models.

Scope and next steps. Our evaluation uses 3T public datasets, so direct claims about low-field or portable MRI are premature. Future work should evaluate MR-DiffRecon on representative low-field scanners and constrained deployment hardware, including latency, memory, energy use, and prospective diagnostic performance.

Table 3: Efficiency comparison across MRI reconstruction methods. MR-DiffRecon achieves the highest SSIM and PSNR, the lowest parameter count overall, and the lowest FLOPs among the compared all-in-one methods. Accuracy is reported on Cine SAX at $\times 10$; FLOPs are measured per output slice in the CMRxRecon 10-coil setting.

Method	Params(M) $\downarrow$	FLOPs(G) $\downarrow$	SSIM $\uparrow$	PSNR $\uparrow$	NMSE($\times 10^{-2}$) $\downarrow$
E2E-VarNet [15]	93.08	370.14	0.9744	42.05	1.6
HUMUS-Net-L [4]	160.97	575.16	0.9791	42.96	1.3
PromptIR [14]	284.74	4455.11	0.9659	40.16	2.5
PromptMR [19]	69.71	1784.75	0.9865	45.58	1.1
MR-DiffRecon	53.33	1647.04	0.9947	46.63	0.7

Table 4: Cumulative component ablation on Cine SAX ($\times 10$).

Dual Attention	Single-Scale Align.	Multi-Scale Align.	Reverse Diffusion	SSIM $\uparrow$	PSNR $\uparrow$	NMSE($\times 10^{-2}$) $\downarrow$
$\checkmark$	$\times$	$\times$	$\times$	0.9879	45.93	1.1
$\checkmark$	$\checkmark$	$\times$	$\times$	0.9885	46.06	1.1
$\checkmark$	$\checkmark$	$\checkmark$	$\times$	0.9891	46.18	1.0
$\checkmark$	$\checkmark$	$\checkmark$	$\checkmark$	0.9947	46.63	0.7

7 Conclusion

We presented MR-DiffRecon, an efficient all-in-one MRI reconstruction framework that combines multi-scale adjacent-slice alignment, implicit prompting, and training-time reverse diffusion within an 8-cascade unrolled backbone. Within each dataset, a single model handles multiple acquisition settings, while the same architectural design is validated on cardiac and knee MRI using dataset-specific training. MR-DiffRecon improves SSIM and PSNR over the evaluated baselines while reducing parameters and inference FLOPs relative to PromptMR, and by removing diffusion sampling at deployment it improves the accuracy–efficiency trade-off at no added inference cost. Future work will target prospective, low-field, and resource-constrained deployment, including on-device latency, memory, and energy evaluation.

References

1. Dai, J., Qi, H., Xiong, Y., Li, Y., Zhang, G., Hu, H., Wei, Y.: Deformable convolutional networks (2017), <https://arxiv.org/abs/1703.06211>
2. Donoho, D.L.: Compressed sensing. IEEE Transactions on information theory **52**(4), 1289–1306 (2006)
3. Dudhane, A., Zamir, S.W., Khan, S., Khan, F.S., Yang, M.H.: Burst image restoration and enhancement. In: CVPR (2022)
4. Fabian, Z., Tinaz, B., Soltanolkotabi, M.: HUMUS-Net: Hybrid unrolled multi-scale network architecture for accelerated MRI reconstruction. Advances in Neural Information Processing Systems **35**, 25306–25319 (2022)

5. Geethanath, S., Vaughan Jr, J.T.: Accessible magnetic resonance imaging: A review. *Journal of Magnetic Resonance Imaging* **49**(7), e65–e77 (2019). <https://doi.org/10.1002/jmri.26638>
6. Güngör, A., Dar, S.U.H., Öztürk, Ş., Korkmaz, Y., Bedel, H.A., Elmas, G., Ozbey, M., Çukur, T.: Adaptive diffusion priors for accelerated MRI reconstruction. *Medical Image Analysis* **88**, 102872 (Aug 2023). <https://doi.org/10.1016/j.media.2023.102872>, <https://doi.org/10.1016/j.media.2023.102872>
7. Guo, P., Mei, Y., Zhou, J., Jiang, S., Patel, V.M.: Reconformer: Accelerated mri reconstruction using recurrent transformer. *arXiv preprint arXiv:2201.09376* (2022)
8. Hammernik, K., Klatzer, T., Kobler, E., Recht, M.P., Sodickson, D.K., Pock, T., Knoll, F.: Learning a variational network for reconstruction of accelerated MRI data. *Magnetic Resonance in Medicine* **79**(6), 3055–3071 (2018). <https://doi.org/10.1002/mrm.26977>
9. Kingma, D.P., Ba, J.: Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980* (2014)
10. Korkmaz, Y., Cukur, T., Patel, V.: Self-supervised mri reconstruction with unrolled diffusion models (2023)
11. Lustig, M., Donoho, D., Santos, J.M., Pauly, J.M.: Sparse mri: The application of compressed sensing for rapid mr imaging. *Magnetic resonance in medicine* **58**(6), 1182–1195 (2007)
12. Mao, Y., Jiang, L., Chen, X., Li, C.: Disc-diff: Disentangled conditional diffusion model for multi-contrast mri super-resolution (2023)
13. Ogbole, G.I., Adeyomoye, A.O., Badu-Pepurah, A., Mensah, Y., Nzeh, D.A.: Survey of magnetic resonance imaging availability in west africa. *Pan African Medical Journal* **30**, 240 (2018). <https://doi.org/10.11604/pamj.2018.30.240.14000>
14. Potlapalli, V., Zamir, S.W., Khan, S., Khan, F.: Promptir: Prompting for all-in-one image restoration. In: *Thirty-seventh Conference on Neural Information Processing Systems* (2023)
15. Sriram, A., Zbontar, J., Murrell, T., Defazio, A., Zitnick, C.L., Yakubova, N., Knoll, F., Johnson, P.: End-to-end variational networks for accelerated mri reconstruction. In: *Medical Image Computing and Computer Assisted Intervention–MICCAI 2020: 23rd International Conference, Lima, Peru, October 4–8, 2020, Proceedings, Part II* 23. pp. 64–73. Springer (2020)
16. Wang, C., Lyu, J., Wang, S., Qin, C., Guo, K., Zhang, X., Yu, X., Li, Y., Wang, F., Jin, J., et al.: Cmrrecon: An open cardiac mri dataset for the competition of accelerated image reconstruction. *arXiv preprint arXiv:2309.10836* (2023)
17. Wang, X., Wang, S., Xiong, H., Xuan, K., Zhuang, Z., Liu, M., Shen, Z., Zhao, X., Zhang, L., Wang, Q.: Spatial attention-based implicit neural representation for arbitrary reduction of mri slice spacing (2023)
18. Xiang, T., Yue, W., Lin, Y., Yang, J., Wang, Z., Li, X.: Diffcmr: Fast cardiac mri reconstruction with diffusion probabilistic models (2023)
19. Xin, B., Ye, M., Axel, L., Metaxas, D.N.: Fill the k-space and refine the image: Prompting for dynamic and multi-contrast mri reconstruction. *arXiv preprint arXiv:2309.13839* (2023)
20. Yang, H., Li, J., Lui, L.M., Ying, S., Shi, J., Zeng, T.: Fast mri reconstruction via edge attention. *Communications in Computational Physics* **33**(5), 1409–1431 (2023). <https://doi.org/10.4208/cicp.oa-2022-0309>, <http://dx.doi.org/10.4208/cicp.OA-2022-0309>
21. Zbontar, J., Knoll, F., Sriram, A., Murrell, T., Huang, Z., Muckley, M.J., Defazio, A., Stern, R., Johnson, P., Bruno, M., Parente, M., Geras, K.J., Katsnelson, J.,

- Chandarana, H., Zhang, Z., Drozdal, M., Romero, A., Rabbat, M., Vincent, P., Yakubova, N., Pinkerton, J., Wang, D., Owens, E., Zitnick, C.L., Recht, M.P., Sodickson, D.K., Lui, Y.W.: fastMRI: An open dataset and benchmarks for accelerated MRI. arXiv preprint arXiv:1811.08839 (2018)
22. Zhang, J., Chi, Y., Lyu, J., Yang, W., Tian, Y.: Dual arbitrary scale super-resolution for multi-contrast mri (2023)