

FieldRF: Anatomy-Conditioned 3D Rectified Flow for Unified Cross-Field MRI Harmonization

Haiyu Song¹, Lei Xiang^{1,2}, Zongjian Yang¹, Yang Wu¹, Dinggang Shen^{1,2,3},
and Qian Wang^{1,3*}

¹ School of Biomedical Engineering & State Key Laboratory of Advanced Medical Materials and Devices, ShanghaiTech University, Shanghai 201210, China

wangqian2@shanghaitech.edu.cn

² Shanghai United Imaging Intelligence Co., Ltd., Shanghai 200230, China

³ Shanghai Clinical Research and Trial Center, Shanghai 201210, China

Abstract. Cross-field MRI harmonization must change acquisition appearance while preserving anatomy. Yet paired scans are scarce, and pairwise translation scales poorly. We formulate this task as anatomy-conditioned volumetric transport and present FieldRF, one 3D rectified-flow model shared across modality-field domains. FieldRF separates conditioning into subject anatomy, which anchors structure, and a modality-field label, which controls target appearance. We leverage complementary pretrained medical foundation models: anatomical foundation representations provide field-robust structural guidance, while volumetric generative models make 384³ full-volume latent transport tractable. Training first learns conditional flow from unpaired scans without adversarial supervision; limited paired scans can optionally refine cross-field transport. One model covers all within-modality field directions; official evaluation shows competitive performance across T1W, T2W, and T2FLAIR.

Keywords: MRI harmonization · Unpaired learning · Rectified flow · Semantic conditioning · Latent generative model

1 Introduction

Magnetic field strength affects MRI signal-to-noise ratio, contrast, and visible anatomical detail. Scans acquired at 0.1T, 1.5T, 3T, 5T, and 7T therefore cannot be pooled as if they came from a common distribution. Cross-field harmonization seeks to transform an image to a requested target-field appearance while preserving subject-specific anatomy.

The MRIFields2026 challenge [11] makes this setting particularly demanding. Its retrospective training data are unpaired, and only a small prospective paired training cohort is available before evaluation on a separate validation set. Moreover, the any-to-any task requires a unified model rather than a collection of direction-specific generators. Conventional unpaired approaches such as CycleGAN [19] and CUT [12] learn mappings between image domains, but scaling

* Corresponding author.

pairwise formulations to many field strengths multiplies the number of models. The challenge therefore calls for a unified volumetric alternative to pairwise or adversarial translation baselines.

FieldRF formulates harmonization as anatomy- and domain-conditioned 3D rectified flow. BrainID provides an explicit subject-anatomy condition, the domain label specifies the requested modality and field, and one shared volumetric velocity field learns all domains. Separating anatomy and domain conditions decouples subject-specific structure from acquisition-dependent appearance. This differs from the direct pair-specific mappings of CycleGAN and CUT and the adversarial multi-domain translation of StarGANv2 [3]. The model requires no adversarial realism supervision, although this cannot rule out hallucinated content.

Rather than training the volumetric system entirely from scratch, FieldRF reuses pretrained medical models. Frozen BrainID [9] provides spatial anatomy through ControlNet [16]; its features have 0.9776 cross-field cosine similarity. The MAISI VAE [6] compresses 384^3 volumes to 4×96^3 latents, MAISI-v2 [18] initializes the 3D flow backbone, and only the VAE decoder is adapted before freezing.

Our contributions are:

1. We reformulate cross-field MRI harmonization from pairwise image translation to anatomy-conditioned volumetric latent transport, learning a continuous trajectory that preserves subject-specific anatomy while transporting acquisition-dependent appearance to the requested domain.
2. We introduce a factorized conditioning strategy that separates a spatial subject-anatomy condition from an explicit target modality-field appearance condition within one shared 3D velocity field. Complementary pretrained anatomical and volumetric generative foundation models enable field-robust structural guidance and scalable full-volume 3D latent transport.
3. We develop an unpaired-first, data-efficient training strategy that learns domain-conditioned transport from retrospective unpaired volumes without adversarial supervision and optionally exploits limited paired scans to refine cross-field transport and latent endpoint alignment.

2 Related Work

MRI harmonization. DeepHarmony learns paired scanner-to-scanner contrast correction [5], while CALAMITI disentangles anatomy and contrast for unpaired multi-site harmonization [20]. Harmonizing Flows aligns images to a learned normalizing-flow distribution [1]. More recently, HCLD performs unpaired 3D brain MRI harmonization with conditional latent diffusion and explicit content/style constraints [15]. FieldRF instead takes a transport-oriented approach, learning $T_\theta(z_s \rightarrow z_t \mid b_s, d_t)$ to map the encoded source through a subject-specific endpoint to the requested domain, rather than synthesizing from a conditional generative distribution.

Unpaired translation and generative transport. CycleGAN introduces cycle consistency for unpaired translation [19], whereas CUT maximizes patchwise mutual information between input and output [12]; StarGANv2 addresses translation among multiple domains [3]. Diffusion and neural Schrödinger bridges formulate distribution transport through a stochastic bridge [4, 7]; FieldRF uses one source-conditioned 3D rectified flow for all requested domains.

Latent flow, spatial control, and anatomy. Latent diffusion moves generative modeling to a compressed representation [13]; Flow Matching [8] and Rectified Flow [10] learn continuous transport fields. MAISI and MAISI-v2 respectively adapt latent diffusion and rectified flow to 3D medical images [6, 18]. ControlNet injects spatial conditions through residual branches [16], and BrainID learns contrast-agnostic anatomical representations [9]. We freeze BrainID as the conditioner; the complete anatomy-conditioned rectified-flow pipeline, rather than BrainID alone, performs harmonization.

3 Method

3.1 Problem Formulation

Let x_s denote an MRI from source domain d_s , and d_t the requested target domain. A domain is the joint modality-field combination. The base model uses individual retrospective observations (x, d) without paired (x_s, x_t) scans; an optional second phase fine-tunes cross-field transport using the small prospective paired set. Let E denote the frozen VAE encoder followed by its fixed latent scaling, and D the corresponding inverse-scaled decoder, adapted alone before being frozen. Let B_{ID} be the pretrained frozen BrainID conditioner and I_θ, G_θ the inversion and generation operators induced by one shared velocity field v_θ . We seek

$$e_s = I_\theta(E(x_s); B_{\text{ID}}(x_s), d_s), \quad \hat{x}_{s \rightarrow t} = D(G_\theta(e_s; B_{\text{ID}}(x_s), d_t)), \quad (1)$$

Our central modeling assumption is that $B_{\text{ID}}(x)$ preserves anatomy while suppressing most acquisition appearance. Each training image supervises the shared velocity field under its observed domain label,

$$(B_{\text{ID}}(x), d) \mapsto E(x), \quad (2)$$

Cross-field inference then combines the source inversion endpoint and source anatomy with $d_t \neq d_s$. The paired objective below directly supervises this source-to-target route.

3.2 Latent Rectified Flow

Let $z_0 = E(x)$ be the scaled target latent and $\epsilon \sim \mathcal{N}(0, I)$. For $t \sim \mathcal{U}(0, T)$, we define

$$\alpha = t/T, \quad z_t = (1 - \alpha)z_0 + \alpha\epsilon, \quad v^* = z_0 - \epsilon. \quad (3)$$

Generation is integrated from the endpoint ($\alpha = 1$) toward the image latent ($\alpha = 0$), so the network predicts the corresponding reverse-time velocity:

$$\hat{v} = v_\theta(z_t, \tau(t), B_{\text{ID}}(x), c_d), \quad \mathcal{L}_{\text{RF}} = \mathbb{E}[\|\hat{v} - v^*\|_1], \quad (4)$$

where $B_{\text{ID}}(x)$ denotes the frozen BrainID spatial feature and c_d is the target-domain label. The resolution-aware timestep transform inherited from MAISI-v2 is

$$\tau(t) = T \frac{r(t/T)}{1 + (r - 1)(t/T)}, \quad r = 1.4 \left(\frac{96^3}{32^3} \right)^{1/3}. \quad (5)$$

During harmonization, target-domain reverse generation uses 20 Euler steps. FieldRF-U uses four source-inversion steps, whereas FieldRF-FT uses 20.

3.3 Semantic Control and Domain Conditioning

The backbone is a four-level 3D U-Net initialized from the released MAISI MR-brain rectified-flow checkpoint. Its main input is the four-channel noisy latent. A randomly initialized MONAI ControlNet receives the same latent and a frozen BrainID spatial feature. Its conditioning embedding projects the anatomy feature to the latent grid before the control residuals are injected into the U-Net down blocks and middle block.

Modality and field strength are represented jointly:

$$c_d = 5c_{\text{modality}} + c_{\text{field}}, \quad c_d \in \{0, \dots, 14\}. \quad (6)$$

All convolutional and attention parameters are shared across labels. The model therefore contains no field-pair-specific generators or decoders.

3.4 Paired RF Fine-tuning

The paired stage starts from the unpaired FieldRF model. Let (x_s, x_t) be a same-subject, same-modality cross-field pair. Denote its latents by z_s, z_t , BrainID features by b_s, b_t , and inversion endpoints by e_s, e_t . A residual endpoint adapter A_ϕ gives $e'_s = A_\phi(e_s, d_s)$ and $e'_t = A_\phi(e_t, d_t)$. At a common sampled time, we construct RF paths from both endpoints to the real target latent and use

$$\begin{aligned} \mathcal{L}_{\text{cross}} &= \|v_\theta(z_\tau^{s \rightarrow t}, \tau, b_s, d_t) - (z_t - e'_s)\|_1, \\ \mathcal{L}_{\text{id}} &= \|v_\theta(z_\tau^{t \rightarrow t}, \tau, b_t, d_t) - (z_t - e'_t)\|_1. \end{aligned} \quad (7)$$

Every fourth update, a 20-step source-to-target rollout $\tilde{z}_t = G_\theta(e'_s; b_s, d_t)$ defines the latent endpoint consistency

$$\mathcal{L}_{\text{latent}} = \|\tilde{z}_t - z_t\|_1. \quad (8)$$

During optimization it is augmented by a 0.02-weighted latent edge term and a 0.25-weighted SmoothL1 term against a stop-gradient target-identity rollout;

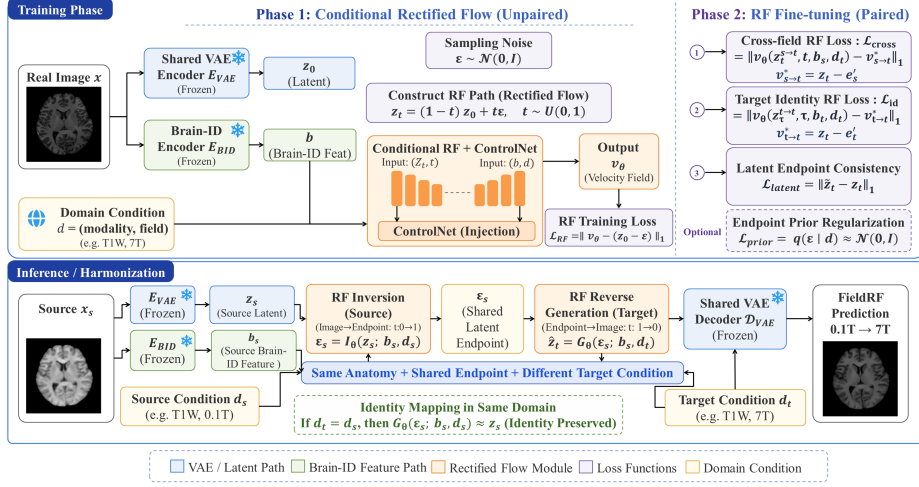


Fig. 1. Overview of the proposed FieldRF pipeline. Phase 1 learns conditional rectified flow from unpaired real images; Phase 2 can fine-tune the same RF using paired cross-field scans. At inference, source inversion and target generation share the frozen BrainID condition while changing only the domain label. The dashed optional endpoint prior is not used in the reported experiments.

their sum is denoted $\mathcal{L}_{\text{latent}}^{\text{impl}}$. The paired objective is $\mathcal{L}_{\text{FT}} = \mathcal{L}_{\text{cross}} + 0.25\mathcal{L}_{\text{id}} + \mathcal{L}_{\text{latent}}^{\text{impl}}$ on rollout updates and omits the last term otherwise. BrainID and the VAE stay frozen; the U-Net, ControlNet, and endpoint adapter are updated without an image-space or adversarial loss. The adapter is initialized from an unpaired pre-alignment checkpoint; during paired fine-tuning it receives gradients only through the two local RF losses.

4 Experiments

4.1 Data and Evaluation Protocol

We use the MRIFields2026 release dated 14 April 2026 [11]: T1W, T2W, and T2FLAIR MRI at 0.1T, 1.5T, 3T, 5T, and 7T. The 1,939 retrospective volumes, listed in this field order, comprise T1W 100/104/143/121/235 (703), T2W 100/221/142/43/81 (587), and T2FLAIR 99/221/143/102/84 (649); all are included without subject-level exclusions. Prospective training and validation each provide three subjects per modality–field combination. Retrospective scans train base FieldRF; paired prospective training scans fine-tune the base-lines and FieldRF-FT, while validation cases never update parameters. Volumes are resampled to 384^3 ; the VAE operates on 4×96^3 latents, and BrainID supplies the frozen spatial anatomy condition.

We report SSIM [14], normalized RMSE, LPIPS [17], Dice, and volume similarity on the prospective validation data. The structural metrics use SynthSeg [2].

The T1W comparison macro-averages the four source fields. The challenge evaluates the axial slab $z = [150, 180]$; all metrics in Table 1 are computed on these same 30 slices. We additionally report the official Task 3 server evaluation, which covers all 20 ordered non-identity mappings for each modality (180 cases in total) and subsumes the Task 2 directions.

4.2 Baselines

CycleGAN. CycleGAN [19] uses two generators, two discriminators, and cycle consistency. We train twelve pair-specific models (four fields and three modalities), each with nine-block, 64-filter ResNet generators, three-layer PatchGAN discriminators, LSGAN, cycle weight 10, and identity ratio 0.5. Retrospective pretraining uses 2×10^{-4} for 50 epochs plus 10 epochs of linear decay; the generator is then fine-tuned for 50 epochs on paired prospective scans with $\mathcal{L}_1 + 0.1\mathcal{L}_{\text{LPIPS}}$.

CUT. CUT [12] uses one-sided adversarial translation with multilayer PatchNCE. Each of its twelve pair-specific runs uses one nine-block, 64-filter ResNet generator, one three-layer PatchGAN, LSGAN, and 768 foreground patches from encoder layers 0/4/8/12/16 ($T = 0.07$, $\lambda_{\text{GAN}} = 1$, $\lambda_{\text{NCE}} = 5$). It follows the same 50+10 epoch pretraining and 50-epoch paired fine-tuning schedule as CycleGAN. Each mapping uses the checkpoint with the highest prospective-training SSIM; validation cases are used only for the final server evaluation.

StarGANv2. For the full Task 3 comparison, we use one shared 15-domain StarGANv2 model [3]. Its ResNet generator is pretrained for 25 epochs on retrospective data using adversarial, style, diversity, cycle, and content objectives, then fine-tuned for 50 epochs on paired prospective scans with image, LPIPS, and SSIM losses. We report the fixed epoch-50 model. StarGANv2 trains on $z = [120, 210]$ and is evaluated on $z = [150, 180]$.

All three baselines process axial slices and include paired prospective fine-tuning, whereas the base FieldRF model uses retrospective unpaired data and operates on complete resampled 3D volumes. CycleGAN and CUT use $z = [150, 180]$ for both stages, matching the challenge scoring slab; separate whole-slice pilot runs gave lower validation scores and were not used in Table 1. For the matched Task 3 comparison, we additionally fine-tune the shared FieldRF model on the paired prospective training split.

4.3 Implementation Details

Training uses Adam with learning rate 10^{-4} , five warm-up epochs, cosine decay to 10^{-6} , batch size one, bfloat16 mixed precision, and EMA decay 0.999. Unpaired training lasts 100 epochs. The pretrained MAISI VAE maps 384^3 images to 4×96^3 latents; its encoder stays fixed, and only its decoder is adapted on retrospective MRI before both are frozen. The released MAISI-v2 RF U-Net initializes the backbone, while the 64-channel ControlNet is randomly initialized.

The shared RF backbone is optimized with Eq. (4); adversarial training is disabled. Inference uses the variant-specific inversion schedules stated above and 20 reverse-generation steps.

Paired fine-tuning protocol. We optimize the paired objective for 20 epochs on the prospective training subjects; validation subjects are used only for official server evaluation. We denote the unpaired and paired variants FieldRF-U and FieldRF-FT.

4.4 Main Results

Table 1 gives the cohort-matched T1W comparison. One shared FieldRF-U model obtains the highest SSIM and lowest nRMSE across four directions; it also has the highest volume similarity. CUT has the lowest LPIPS and highest Dice. Thus, FieldRF’s advantage is not uniform across metrics.

Paired fine-tuning changes mean nRMSE, SSIM, and LPIPS from 0.3217, 0.8765, and 0.1313 to 0.2844, 0.8841, and 0.1194. FieldRF leads StarGANv2 in SSIM for every modality (0.0148 mean gain), but trails in nRMSE and LPIPS.

4.5 Diagnostics

BrainID was introduced as a contrast-agnostic anatomical representation and validated through reconstruction and segmentation [9]. Its frozen spatial feature here conditions the shared flow rather than serving as an output metric. Same-modality cross-field cosine similarity is 0.9776 overall: 0.9872 (T1W), 0.9734 (T2W), and 0.9718 (T2FLAIR). This tests condition robustness, not output preservation; Table 1 and Fig. 2 give complementary output-level evidence.

Figure 2 uses submitted-system NIFTI predictions and 7T targets, with identical reorientation/rotation and a common error scale. Table 1 retains all subjects; Table 2 covers all modalities and ordered field directions.

5 Discussion and Conclusion

FieldRF targets a unified harmonization regime: one non-adversarial, anatomy-conditioned 3D transport model covers all evaluated within-modality field directions without direction-specific generators.

Pretrained BrainID supplies field-robust anatomy, while the MAISI VAE and MAISI-v2 enable tractable 384^3 latent transport and initialize its 3D backbone. Official evaluation confirms competitive synthesis across all three modalities; paired fine-tuning improves SSIM over the shared StarGANv2 baseline. Velocity regression omits adversarial realism supervision but cannot rule out hallucination. Limitations include VAE-bounded fidelity and the challenge-defined three-case cohort per Task 3 direction.

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

Table 1. Challenge-validation T1W adjusted metrics on twelve cases: three from each 0.1T/1.5T/3T/5T-to-7T direction. Baselines use four pair-specific, paired-fine-tuned models; FieldRF-U uses one shared unpaired model.

Method	Models	SSIM $\uparrow$	nRMSE $\downarrow$	LPIPS $\downarrow$	Dice $\uparrow$	Volume $\uparrow$
CycleGAN [19]	4	0.8995	0.3027	0.0767	0.8261	0.8267
CUT [12]	4	0.8941	0.3416	0.0746	0.8323	0.8168
FieldRF-U	1	0.9081	0.2918	0.1031	0.8049	0.8383

Table 2. Official Task 3 adjusted results over 180 predictions and all 60 directions. Each method uses one shared 15-domain model; FieldRF-FT and StarGANv2 include paired fine-tuning.

Scope	StarGANv2			FieldRF-U			FieldRF-FT		
	nRMSE $\downarrow$	SSIM $\uparrow$	LPIPS $\downarrow$	nRMSE $\downarrow$	SSIM $\uparrow$	LPIPS $\downarrow$	nRMSE $\downarrow$	SSIM $\uparrow$	LPIPS $\downarrow$
T1W	0.2798	0.8782	0.1031	0.3325	0.8788	0.1250	0.3052	0.8842	0.1188
T2W	0.2813	0.8809	0.1048	0.3226	0.8884	0.1218	0.2984	0.8915	0.1083
T2FLAIR	0.2129	0.8487	0.1142	0.3101	0.8623	0.1471	0.2495	0.8765	0.1312
Mean	0.2580	0.8693	0.1073	0.3217	0.8765	0.1313	0.2844	0.8841	0.1194

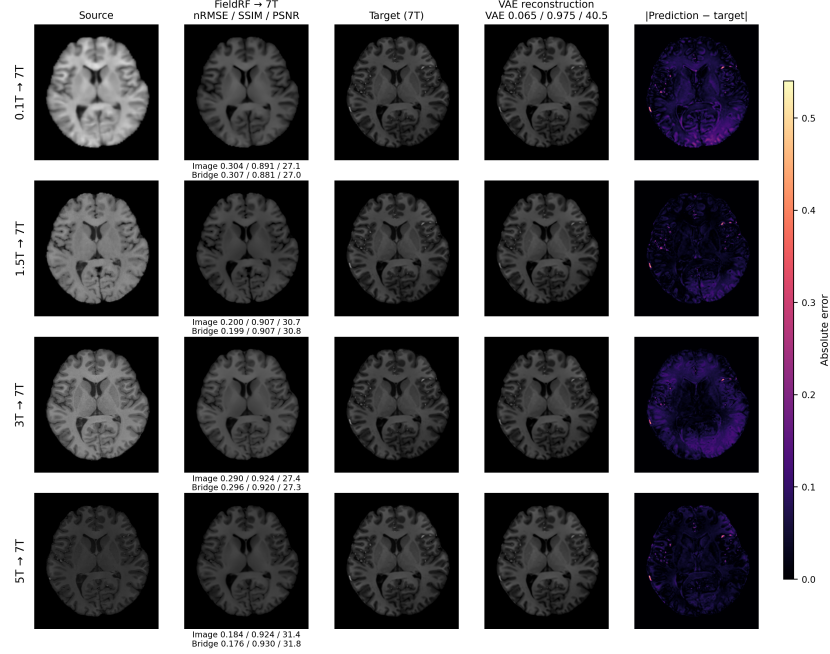


Fig. 2. Illustrative same-subject T1W example (subject 0007): source, FieldRF, target 7T, target VAE reconstruction, and absolute error for four source fields. Triplets compare VAE–target, prediction–target, and prediction–VAE (nRMSE/SSIM/PSNR); errors share one range.

References

1. Beizae, F., Desrosiers, C., Lodygensky, G.A., Dolz, J.: Harmonizing flows: Unsupervised MR harmonization based on normalizing flows. In: *Information Processing in Medical Imaging. Lecture Notes in Computer Science*, vol. 13939, pp. 347–359. Springer (2023). https://doi.org/10.1007/978-3-031-34048-2_27
2. Billot, B., Greve, D.N., Puonti, O., Thielscher, A., Van Leemput, K., Fischl, B., Dalca, A.V., Iglesias, J.E.: SynthSeg: Segmentation of brain MRI scans of any contrast and resolution without retraining. *Medical Image Analysis* **86**, 102789 (2023). <https://doi.org/10.1016/j.media.2023.102789>
3. Choi, Y., Uh, Y., Yoo, J., Ha, J.W.: StarGAN v2: Diverse image synthesis for multiple domains. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 8188–8197 (2020). <https://doi.org/10.1109/CVPR42600.2020.00821>
4. De Bortoli, V., Thornton, J., Heng, J., Doucet, A.: Diffusion schrödinger bridge with applications to score-based generative modeling. In: *Advances in Neural Information Processing Systems*. vol. 34, pp. 17695–17709 (2021)
5. Dewey, B.E., Zhao, C., Reinhold, J.C., Carass, A., Fitzgerald, K.C., Sotirchos, E.S., Saidha, S., Oh, J., Pham, D.L., Calabresi, P.A., van Zijl, P.C.M., Prince, J.L.: DeepHarmony: A deep learning approach to contrast harmonization across scanner changes. *Magnetic Resonance Imaging* **64**, 160–170 (2019). <https://doi.org/10.1016/j.mri.2019.05.041>
6. Guo, P., Zhao, C., Yang, D., Xu, Z., Nath, V., Tang, Y., Simon, B., Belue, M., Harmon, S., Turkbey, B., Xu, D.: MAISI: Medical AI for synthetic imaging. In: *Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision*. pp. 4430–4441 (2025). <https://doi.org/10.1109/WACV61041.2025.00435>
7. Kim, B., Kwon, G., Kim, K., Ye, J.C.: Unpaired image-to-image translation via neural schrödinger bridge. In: *International Conference on Learning Representations* (2024), https://proceedings.iclr.cc/paper_files/paper/2024/hash/5491280797f3192b895bce84eb83df8d-Abstract-Conference.html
8. Lipman, Y., Chen, R.T.Q., Ben-Hamu, H., Nickel, M., Le, M.: Flow matching for generative modeling. In: *International Conference on Learning Representations* (2023), <https://openreview.net/forum?id=PqvMRDCJT9t>
9. Liu, P., Puonti, O., Hu, X., Alexander, D.C., Iglesias, J.E.: Brain-ID: Learning contrast-agnostic anatomical representations for brain imaging. In: *Computer Vision – ECCV 2024*. pp. 322–340. Springer (2024). https://doi.org/10.1007/978-3-031-73254-6_19
10. Liu, X., Gong, C., Liu, Q.: Flow straight and fast: Learning to generate and transfer data with rectified flow. In: *International Conference on Learning Representations* (2023), <https://openreview.net/forum?id=XVjTT1nw5z>
11. MRIFields2026 Organizers: MRIFields2026: Cross-field mri harmonization challenge. <https://www.synapse.org/Synapse:syn72060672> (2026), challenge website, accessed 22 July 2026
12. Park, T., Efros, A.A., Zhang, R., Zhu, J.Y.: Contrastive learning for unpaired image-to-image translation. In: *Computer Vision – ECCV 2020. Lecture Notes in Computer Science*, vol. 12354, pp. 319–345. Springer (2020). https://doi.org/10.1007/978-3-030-58545-7_19
13. 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). <https://doi.org/10.1109/CVPR52688.2022.01042>

14. Wang, Z., Bovik, A.C., Sheikh, H.R., Simoncelli, E.P.: Image quality assessment: From error visibility to structural similarity. *IEEE Transactions on Image Processing* **13**(4), 600–612 (2004). <https://doi.org/10.1109/TIP.2003.819861>
15. Wu, M., Yu, M., Jing, S., Yap, P.T., Zhang, Z., Liu, M.: Unpaired volumetric harmonization of brain MRI with conditional latent diffusion. *Medical Image Analysis* **107**, 103849 (2026). <https://doi.org/10.1016/j.media.2025.103849>
16. Zhang, L., Rao, A., Agrawala, M.: Adding conditional control to text-to-image diffusion models. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision*. pp. 3836–3847 (2023). <https://doi.org/10.1109/ICCV51070.2023.00355>
17. Zhang, R., Isola, P., Efros, A.A., Shechtman, E., Wang, O.: The unreasonable effectiveness of deep features as a perceptual metric. In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*. pp. 586–595 (2018). <https://doi.org/10.1109/CVPR.2018.00068>
18. Zhao, C., Guo, P., Yang, D., He, Y., Tang, Y., Simon, B., Belue, M., Harmon, S., Turkbey, B., Xu, D.: MAISI-v2: Accelerated 3d high-resolution medical image synthesis with rectified flow and region-specific contrastive loss. *Proceedings of the AAAI Conference on Artificial Intelligence* **40**(15), 13088–13098 (2026). <https://doi.org/10.1609/aaai.v40i15.38309>
19. Zhu, J.Y., Park, T., Isola, P., Efros, A.A.: Unpaired image-to-image translation using cycle-consistent adversarial networks. In: *Proceedings of the IEEE International Conference on Computer Vision*. pp. 2223–2232 (2017). <https://doi.org/10.1109/ICCV.2017.244>
20. Zuo, L., Dewey, B.E., Liu, Y., He, Y., Newsome, S.D., Mowry, E.M., Resnick, S.M., Prince, J.L., Carass, A.: Unsupervised MR harmonization by learning disentangled representations using information bottleneck theory. *NeuroImage* **243**, 118569 (2021). <https://doi.org/10.1016/j.neuroimage.2021.118569>