

# UNSUPERVISED PARALLEL MRI RECONSTRUCTION VIA PROJECTED CONDITIONAL FLOW MATCHING

**Anonymous authors**

Paper under double-blind review

## ABSTRACT

Reconstructing high-quality images from substantially undersampled k-space data for accelerated MRI presents a challenging ill-posed inverse problem. Supervised deep learning has transformed the field by using large amounts of fully sampled ground-truth MR images, either to directly reconstruct undersampled data into fully sampled images with neural networks, or to learn the prior distribution of fully sampled images through generative models. However, in practical scenarios, acquiring ground-truth fully sampled MRI images is not viable due to the inherently slow nature of its data acquisition process. Despite advances in self-supervised/unsupervised MRI reconstruction, the performance remains inadequate at high acceleration rates. To address these gaps, we introduce the Projected Conditional Flow Matching (PCFM) and its unsupervised transformation, which is designed to learn the prior distribution of fully sampled parallel MRI by solely utilizing the undersampled k-space measurements. To reconstruct the image, we establish a novel relationship between the marginal vector field in the measurement space, which generates the associated probability flow in terms of the continuity equation, and the optimal solution to PCFM. This connection results in a cyclic dual-space sampling algorithm for unsupervised reconstruction. Our method was evaluated against contemporary state-of-the-art supervised, self-supervised, and unsupervised baseline techniques on parallel MRI using publicly available datasets fastMRI and CMRxRecon. Experimental results show that our technique significantly surpasses existing self-supervised and unsupervised baselines, while also yielding better performance than most supervised methods. Our code will be available at <https://github.com/anonymous>.

## 1 INTRODUCTION

Magnetic Resonance Imaging (MRI) is a cornerstone of modern medical diagnostics, providing exceptional soft-tissue contrast without the use of ionizing radiation. However, a significant clinical limitation of MRI is its inherently slow data acquisition speed. The time required for a scan is directly proportional to the amount of data that must be acquired in the MRI raw data space, known as k-space. The development of multi-coil receiver arrays, replacing single coils with smaller, localized ones (Roemer et al., 1990; Sodickson & Manning, 1997), enabled accelerated data acquisition through parallel imaging. Key algorithms, SENSitivity Encoding (SENSE) (Pruessmann et al., 1999) and GeneRalized Autocalibrating Partially Parallel Acquisitions (GRAPPA) (Griswold et al., 2002), laid the groundwork for this. By combining parallel imaging with sparsity-promoting terms, later compressed sensing (CS)-based methods achieve higher acceleration than traditional techniques (Lustig et al., 2007; 2008).

Building on this concept, data-driven models, particularly physics-informed “unrolled” neural networks that emulate classical iterative optimization methods and leverage prior knowledge encoded by convolutional neural networks (CNNs) (Ulyanov et al., 2018), attain cutting-edge outcomes in terms of both reconstruction quality and speed (Aggarwal et al., 2018; Hammernik et al., 2018). More recently, advancements in accelerated MRI reconstruction have been pursued through modern generative models. These generative models, rather than focusing on learning a single point estimate, are designed to approximate the entire probability distribution of high-quality MR images (Mardani et al., 2018; Tezcan et al., 2018; Song et al., 2022; Chung & Ye, 2022). This approach aids in solving the inverse problem by enabling sampling from the posterior distribution  $p(\mathbf{x} | \mathbf{y})$ ,

where  $\mathbf{x}$  represents the target fully sampled image, and  $\mathbf{y}$  includes the undersampled multi-coil k-space measurement. Despite achieving high reconstruction accuracy, these supervised frameworks necessitate access to vast collections of fully sampled ground-truth images for training, which are not only costly but also commonly unattainable. Emerging self-supervised techniques aim to reduce the reliance on fully sampled MRI datasets during training (Wang et al., 2025). Nevertheless, they fall short in accuracy when dealing with highly undersampled data, e.g.,  $8\times$  accelerated MRI.

To address these challenges, we introduce an unsupervised generative model for parallel MRI reconstruction requiring solely undersampled k-space data for training. Drawing inspiration from the generalized Stein’s Unbiased Risk Estimator (Stein, 1981; Eldar, 2008), we present the projected conditional flow matching (PCFM) objective alongside its unsupervised adaptation, which facilitates learning the prior distribution of fully sampled MRI using only undersampled data. Given that a closed-form solution to the projection operator is intractable for parallel MRI, we propose to employ a numerical method to approximate the projection operator during training. Subsequently, we derive a new connection between the probability flow in the measurement space under projection and the optimal solution to the proposed PCFM objective, introducing a reconstruction algorithm based on the optimal PCFM solution. The proposed approach is also capable of dealing with noisy measurements under the formulation. Our framework is evaluated using two public parallel brain and cardiac MRI datasets, demonstrating superior reconstruction performance over existing baselines, even when trained solely with undersampled k-space measurements.

## 2 BACKGROUND

### 2.1 PARALLEL MRI RECONSTRUCTION

The fundamental principle of parallel MRI builds on the fact that each coil in an array receives a distinct spatial sensitivity profile, that is, a unique spatially weighted view of the underlying anatomy. This spatial encoding ability can be used to compensate for the spatial information lost when k-space is undersampled. Formally, the forward model of parallel (also known as multi-coil) MRI writes

$$\mathbf{y}_s = \mathbf{A}_s \mathbf{x} + \mathbf{e}, \quad (1)$$

where  $s$  indexes the randomness in the undersampling mask,  $\mathbf{x} \in \mathcal{X} \subset \mathbb{C}^D$  is the underlying fully sampled complex-valued image,  $\mathbf{y}_s = [\mathbf{y}_{s,1}^\top, \dots, \mathbf{y}_{s,C}^\top]^\top \in \mathcal{Y} \subset \mathbb{C}^{Cd}$  is the acquired k-space measurements from  $C$  receiver coils with  $d \leq D$ , and  $\mathbf{e} \sim \mathcal{CN}(\mathbf{0}, \sigma_0^2 \mathbf{I}_{Cd})$  denotes measurement noise. In particular, the *coil-combined* forward operator  $\mathbf{A}_s$  is defined as

$$\mathbf{A}_s \triangleq \begin{bmatrix} \mathbf{M}_s \mathbf{F} \mathbf{S}_1 \\ \vdots \\ \mathbf{M}_s \mathbf{F} \mathbf{S}_C \end{bmatrix} \in \mathbb{C}^{Cd \times D}, \quad (2)$$

which is composed of the undersampling mask  $\mathbf{M}_s \in \{0, 1\}^{d \times D}$ , the discrete Fourier transform  $\mathbf{F} \in \mathbb{C}^{D \times D}$ , and diagonal matrices  $\mathbf{S}_c \in \mathbb{C}^{D \times D}$  representing the sensitivity maps. For parallel MRI with acceleration factor  $\alpha \triangleq D/d > 1$ , Eq. (2) can be rank-deficient with a non-zero null space, leading to a challenging ill-posed inverse problem. For simplicity, we assume that the randomness in  $s$  is incorporated in  $\mathbf{y}$  in the following, and thus denote the forward operator as  $\mathbf{A}$ .

### 2.2 CONDITIONAL FLOW MATCHING

Conditional flow matching (CFM) (Lipman et al., 2023) provides a simulation-free technique to learn a continuous normalizing flow (Chen et al., 2018) that transforms a base distribution  $p_1$  to a target distribution  $p_0$ . In this work, we assume  $p_1^X = \mathcal{CN}(\mathbf{0}, 2\mathbf{I}_D)$ , and  $p_0^X$  produces the underlying fully sampled images, where the superscript  $X$  indicates that the flow is in the fully sampled image space  $\mathcal{X}$ . This transformation can be specified by an ordinary differential equation (ODE) with a time-dependent smooth vector field  $\mathbf{u}_t^X: [0, 1] \times \mathbb{C}^D \rightarrow \mathbb{C}^D$ , i.e.,  $d\mathbf{x}_t = \mathbf{u}_t^X(\mathbf{x}_t)dt$ . This induces a probability path  $p_t^X$  as the push-forward distribution of  $p_0^X$  by the ODE dynamics, satisfying the continuity equation (Villani et al., 2009):

$$\frac{\partial p_t^X}{\partial t} + \nabla \cdot (p_t^X \mathbf{u}_t^X) = 0. \quad (3)$$

CFM proposes to construct such a marginal vector field  $\mathbf{u}_t^X$  by introducing the conditional variable  $\mathbf{z}^X \sim q(\mathbf{z}^X)$  and conditional vector field  $\mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X)$ . In this paper, we consider the popular choice of  $\mathbf{z}^X = (\mathbf{x}_0, \mathbf{x}_1)$  and the independent coupling  $q = p_0^X \times p_1^X$  as generalized by (Tong et al., 2023). Then, we assume a conditional probability path  $p_t^X(\cdot | \mathbf{z}^X)$  that satisfies the boundary conditions  $p_0^X = \mathbb{E}_{q(\mathbf{z}^X)} [p_0^X(\cdot | \mathbf{z}^X)]$  and  $p_1^X = \mathbb{E}_{q(\mathbf{z}^X)} [p_1^X(\cdot | \mathbf{z}^X)]$ . One example is the conditional optimal transport (OT) path  $p_t^X(\mathbf{x} | \mathbf{z}^X) = \delta_{a_t \mathbf{x}_0 + b_t \mathbf{x}_1}(\mathbf{x})$  with  $a_t = 1 - t$  and  $b_t = t$  (Lipman et al., 2023; Liu et al., 2023). This leads to the conditional vector field  $\mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X) = a'_t \mathbf{x}_0 + b'_t \mathbf{x}_1$  by the continuity equation w.r.t.  $p_t^X(\mathbf{x} | \mathbf{z}^X)$ , where  $a'_t \triangleq \frac{da_t}{dt}$  and  $b'_t \triangleq \frac{db_t}{dt}$ . Then by verifying Eq. (3), we can show that the marginal vector field

$$\mathbf{u}_t^X(\mathbf{x}) \triangleq \mathbb{E}_{q(\mathbf{z}^X)} \left[ \mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X) \frac{p_t^X(\mathbf{x} | \mathbf{z}^X)}{p_t^X(\mathbf{x})} \right] \quad (4)$$

generates the probability flow  $p_t^X$ . Meanwhile, learning of the flow is achieved by minimizing the conditional flow matching objective:

$$\mathcal{L}_{\text{CFM}}(\theta) \triangleq \mathbb{E}_{t, q(\mathbf{z}^X), p_t^X(\mathbf{x} | \mathbf{z}^X)} \|\mathbf{h}_\theta^X(\mathbf{x}, t) - \mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X)\|_2^2, \quad (5)$$

where  $\mathbf{h}_\theta^X(\cdot, t)$  is a network that predicts the  $\mathcal{X}$ -space marginal vector field.

### 3 METHOD

This section introduces the proposed framework designed to reconstruct parallel MRI using only undersampled k-space data. The framework integrates two principal components: (1) the projected conditional flow matching (PCFM) objective along with its unsupervised transformation as a new formulation for learning the prior (Section 3.1), and (2) a new reconstruction algorithm for inference that exploits the relationship between measurement-space probability flow and the optimal solution to PCFM (Section 3.2).

#### 3.1 PROJECTED CONDITIONAL FLOW MATCHING

Due to the scarcity of fully sampled ground-truth signal  $\mathbf{x}_0$ , optimizing the  $\mathcal{X}$ -space CFM objective from Eq. (5) using a large dataset of fully sampled MRI scans is impractical because the conditional path and vector field within the  $\mathcal{X}$ -space are not accessible. To address this, we introduce a  $\mathcal{Y}$ -space conditional probability path  $p_t^Y(\mathbf{y} | \mathbf{z}^Y) = \delta_{a_t \mathbf{y}_0 + b_t \mathbf{y}_1}(\mathbf{y})$ , where  $\mathbf{z}^Y \triangleq (\mathbf{y}_0, \mathbf{y}_1)$ , in which  $\mathbf{y}_0 = \mathbf{A}\mathbf{x}_0 + \mathbf{e}_0$  is the undersampled k-space measurements from parallel MRI, and  $\mathbf{y}_1 = \mathbf{A}\mathbf{x}_1 + \mathbf{e}_1$  with  $\mathbf{e}_1 \sim \mathcal{CN}(\mathbf{0}, \sigma_1^2 \mathbf{I}_{Cd})$  is the projected noise sampled from the base distribution  $p_1^X$ . This leads to

$$\begin{aligned} p_t^Y(\mathbf{y} | \mathbf{z}^X) &= \int p_t^Y(\mathbf{y} | \mathbf{z}^Y) p(\mathbf{z}^Y | \mathbf{z}^X) d\mathbf{z}^Y \\ &= \mathcal{CN}(\mathbf{y} | \mathbf{A}(a_t \mathbf{x}_0 + b_t \mathbf{x}_1), (a_t^2 \sigma_0^2 + b_t^2 \sigma_1^2) \mathbf{I}_{Cd}). \end{aligned} \quad (6)$$

Fig. 1 depicts the graphical model of the random variables.

Since  $\mathbf{x}_0$  and  $\mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X) = a'_t \mathbf{x}_0 + b'_t \mathbf{x}_1$  are unknown, and the forward operator  $\mathbf{A}$  is rank-deficient, we can only expect to optimize the CFM objective Eq. (5) in the range space  $\mathcal{R}(\mathbf{A}^\top)$  of  $\mathbf{A}^\top$  (Eldar, 2008). Therefore, we propose to optimize the following objective function

$$\mathcal{L}_{\text{PCFM}}(\theta) \triangleq \mathbb{E}_{t, q(\mathbf{z}^X), p_t^X(\mathbf{x} | \mathbf{z}^X), p_t^Y(\mathbf{y} | \mathbf{z}^X)} \| \mathbf{P}[\mathbf{v}_\theta^X(\mathbf{y}, t) - \mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X)] \|_2^2, \quad (7)$$

where  $\mathbf{P} \triangleq \mathbf{A}^+ \mathbf{A}$  is the orthogonal projection onto the range space  $\mathcal{R}(\mathbf{A}^\top)$ , with  $\mathbf{A}^+$  denoting the Moore-Penrose pseudoinverse of  $\mathbf{A}$ . Intuitively, this objective projects the CFM error onto the

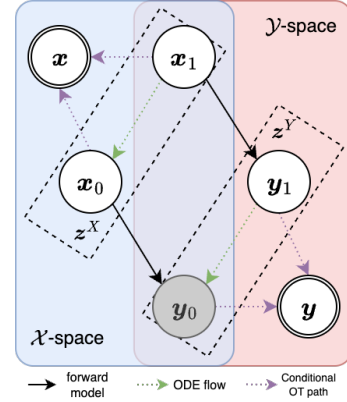


Figure 1: Generation chart of the dual-space conditional probability paths, where observed variables are shaded, and deterministic variables are in double circles. Dotted green arrows indicate deterministic ODE flows, whereas purple ones denote conditional OT paths. The path from  $\mathbf{x}_1$  to  $\mathbf{y}_0$  is traversed either through the  $\mathcal{X}$ -space OR the  $\mathcal{Y}$ -space diagram, as indicated by the background in different colors.

subspace “visible” to the k-space measurements from the mask  $M$ . Thus, we dub this objective function *projected conditional flow matching* (PCFM).

We note that within the framework of the parallel MRI forward model, a closed-form solution for this projection is not available. To tackle this issue, we note the relationships  $P = A^+A = (A^*A)^+A^*A$  and recognize that  $A^*A$  is a positive semi-definite (PSD) matrix. Consequently, we utilize a numerical approximation through the conjugate gradient (CG) method (Appendix C.1), which is suitable for solving the linear system  $A^*A \mathbf{r}_\theta^X = A^*A [\mathbf{v}_\theta^X(\mathbf{y}, t) - \mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X)]$ , where  $\mathbf{r}_\theta^X = P [\mathbf{v}_\theta^X(\mathbf{y}, t) - \mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X)]$ . The optimal solution to the PCFM objective is formalized by the following proposition. The proofs can be found in Appendix A.

**Proposition 1 (Optimal solution to PCFM).** *The minimizer of the PCFM objective is given by*

$$\mathbf{v}_{\theta^*}^X(\mathbf{y}, t) = \mathbb{E}_{q_t(\mathbf{z}^X | \mathbf{y}), p_t^X(\mathbf{x} | \mathbf{z}^X)} [\mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X)] + \mathbf{w}, \quad (8)$$

where  $q_t(\mathbf{z}^X | \mathbf{y}) = \frac{q(\mathbf{z}^X)p_t^Y(\mathbf{y} | \mathbf{z}^X)}{p_t^Y(\mathbf{y})}$ , and  $\mathbf{w}$  is any vector in the null space of  $A$ , i.e.,  $A\mathbf{w} = P\mathbf{w} = \mathbf{0}$ . In particular, when  $\mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X) = a'_t\mathbf{x}_0 + b'_t\mathbf{x}_1$  that is independent of  $\mathbf{x}$ , we have

$$\mathbf{v}_{\theta^*}^X(\mathbf{y}, t) = \mathbb{E}_{q_t(\mathbf{z}^X | \mathbf{y})} [\mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X)] + \mathbf{w}. \quad (9)$$

However, the PCFM objective Eq. (7) still depends on the unknown fully sampled MRI  $\mathbf{x}_0$  through Monte Carlo sampling from  $q(\mathbf{z}^X)$  during training. To address this, inspired by the generalized Stein’s unbiased estimator (Eldar, 2008), we propose to construct an unbiased estimate of the PCFM objective that does not involve the unknown  $\mathbf{x}_0$ . To this end, we notice an induced linear forward model between the dual-space conditional vector fields

$$\mathbf{u}_t^Y(\mathbf{y} | \mathbf{z}^Y) = a'_t\mathbf{y}_0 + b'_t\mathbf{y}_1 = A\mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X) + a'_te_0 + b'_te_1, \quad (10)$$

and the deterministic mapping between the conditional path and the conditional vector field

$$\mathbf{y} = \frac{a_t}{a'_t}\mathbf{u}_t^Y(\mathbf{y} | \mathbf{z}^Y) - b'_t\left(\frac{a_t}{a'_t} - \frac{b_t}{b'_t}\right)\mathbf{y}_1. \quad (11)$$

Based on this, we derive the following unsupervised transformation of the PCFM objective, which does not require fully sampled MRI data for training the vector field predictor.

**Proposition 2 (Unsupervised transformation of PCFM).** *Assuming deterministic conditional probability paths  $\mathbf{x} = a_t\mathbf{x}_0 + b_t\mathbf{x}_1$  and  $\mathbf{y} = a_t\mathbf{y}_0 + b_t\mathbf{y}_1$  with  $\mathbf{y}_0 = A\mathbf{x}_0 + \mathbf{e}_0$  and  $\mathbf{y}_1 = A\mathbf{x}_1 + \mathbf{e}_1$ , where  $\mathbf{e}_0 \sim \mathcal{CN}(\mathbf{0}, \sigma_0^2\mathbf{I}_{Cd})$  and  $\mathbf{e}_1 \sim \mathcal{CN}(\mathbf{0}, \sigma_1^2\mathbf{I}_{Cd})$ , then up to a constant the PCFM objective can be transformed to*

$$\mathbb{E}_{t, q(\mathbf{z}^Y), p_t^Y(\mathbf{y} | \mathbf{z}^Y)} \left[ \left\| P[\mathbf{v}_\theta^X(A^*\mathbf{y}, t) - \hat{\mathbf{u}}_{t, ML}^X] \right\|_2^2 + \frac{2a_t}{a'_t} [(a'_t\sigma_0)^2 + (b'_t\sigma_1)^2] \nabla_{A^*\mathbf{y}} \cdot P\mathbf{v}_\theta^X(A^*\mathbf{y}, t) \right], \quad (12)$$

where  $q(\mathbf{z}^Y) = q(\mathbf{y}_0)q(\mathbf{y}_1) = q(\mathbf{y}_0)\mathbb{E}_{q(\mathbf{x}_1)} [p(\mathbf{y}_1 | \mathbf{x}_1)]$  is sampled by the MRI forward model and Monte Carlo estimation,  $P = A^+A$  is the range-space projection, and

$$\hat{\mathbf{u}}_{t, ML}^X \triangleq (A^*C_t^{-1}A)^+ A^*C_t^{-1}\mathbf{u}_t^Y(\mathbf{y} | \mathbf{z}^Y) \quad (13)$$

with  $C_t = [(a'_t\sigma_0)^2 + (b'_t\sigma_1)^2]\mathbf{I}_d$  is the maximum likelihood solution of the forward model in Eq. (10). Note that  $A^+$  denotes the Moore-Penrose pseudoinverse of  $A$ , which can be approximated by the conjugate gradient method (Appendix C.1).

The network  $\mathbf{v}_\theta^X(\cdot, t)$  takes  $A^*\mathbf{y}$  as input to match the desired dimensionality of the architecture. By optimizing Eq. (12), the optimal solution to PCFM can be determined in an unsupervised learning fashion. The objective can also handle noisy measurements for  $\sigma_0 > 0$ . In the following section, we will delve into reconstructing a fully sampled image utilizing the obtained optimal PCFM solution.

### 3.2 RECONSTRUCTION VIA DUAL-SPACE CYCLIC FLOW INTEGRATION

In the inference phase, to reconstruct the fully sampled image  $\mathbf{x}_0$  given the undersampled measurement  $\mathbf{y}_0$ , we note that given the generative model in Fig. 1, the posterior distribution  $p(\mathbf{x}_0 | \mathbf{y}_0)$  can

be written as

$$\begin{aligned} p(\mathbf{x}_0 | \mathbf{y}_0) &= \int p(\mathbf{x}_0 | \mathbf{x}_1, \mathbf{y}_1, \mathbf{y}_0) p(\mathbf{x}_1 | \mathbf{y}_1, \mathbf{y}_0) p(\mathbf{y}_1 | \mathbf{y}_0) d\mathbf{x}_1 d\mathbf{y}_1 \\ &= \int p(\mathbf{x}_0 | \mathbf{x}_1) p(\mathbf{x}_1 | \mathbf{y}_1) p(\mathbf{y}_1 | \mathbf{y}_0) d\mathbf{x}_1 d\mathbf{y}_1, \end{aligned} \quad (14)$$

which can be evaluated by ancestral Monte-Carlo sampling as shown in Fig. 2. We detail the sampling procedure for each conditional distribution in the following paragraphs.

**Probability vector fields under projection and forward integration.** To sample from the distribution  $p(\mathbf{y}_1 | \mathbf{y}_0)$ , we note that given the conditional probability paths in dual spaces illustrated in Fig. 1, it is feasible to derive a specific relationship between the marginal vector field in the  $\mathcal{Y}$ -space and the optimal solution to the PCFM as discussed in Proposition 1, leading to a delta distribution of  $p(\mathbf{y}_1 | \mathbf{y}_0)$  by the  $\mathcal{Y}$ -space ODE. Recall that the  $\mathcal{Y}$ -space conditional probability path is given by Eq. (6). Denote  $\boldsymbol{\mu}_t(\mathbf{z}^X) \triangleq \mathbf{A}(a_t \mathbf{x}_0 + b_t \mathbf{x}_1)$  and  $\sigma_t^2 = a_t^2 \sigma_0^2 + b_t^2 \sigma_1^2$ . Then the time derivative of  $p_t^Y(\mathbf{y} | \mathbf{z}^X)$  can be written as

$$\begin{aligned} \frac{\partial p_t^Y(\mathbf{y} | \mathbf{z}^X)}{\partial t} &= \frac{\partial p_t^Y(\mathbf{y} | \mathbf{z}^X)}{\partial \boldsymbol{\mu}_t} \cdot \frac{d\boldsymbol{\mu}_t}{dt} + \frac{\partial p_t^Y(\mathbf{y} | \mathbf{z}^X)}{\partial \sigma_t^2} \frac{d\sigma_t^2}{dt} \\ &= -\nabla_{\mathbf{y}} p_t^Y(\mathbf{y} | \mathbf{z}^X) \cdot \mathbf{A}(a'_t \mathbf{x}_0 + b'_t \mathbf{x}_1) + \frac{1}{2} \Delta_{\mathbf{y}} p_t^Y(\mathbf{y} | \mathbf{z}^X) (2a_t a'_t \sigma_0^2 + 2b_t b'_t \sigma_1^2) \\ &= -\nabla_{\mathbf{y}} \cdot (p_t^Y(\mathbf{y} | \mathbf{z}^X) [\mathbf{A} \mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X) - (a_t a'_t \sigma_0^2 + b_t b'_t \sigma_1^2) \nabla_{\mathbf{y}} \log p_t^Y(\mathbf{y} | \mathbf{z}^X)]). \end{aligned} \quad (15)$$

Therefore, by the continuity equation, we know that the  $\mathcal{Y}$ -space conditional vector field

$$\mathbf{u}_t^Y(\mathbf{y} | \mathbf{z}^X) \triangleq \mathbf{A} \mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X) - (a_t a'_t \sigma_0^2 + b_t b'_t \sigma_1^2) \nabla_{\mathbf{y}} \log p_t^Y(\mathbf{y} | \mathbf{z}^X) \quad (16)$$

generates the conditional probability path  $p_t^Y(\mathbf{y} | \mathbf{z}^X)$ . Taking the expectation of Eq. (16) over  $q_t(\mathbf{z}^X | \mathbf{y})$  and leveraging Eq. (9), we can obtain the marginal vector field in the  $\mathcal{Y}$ -space in terms of the optimal PCFM solution and the score function, as described in the following lemma.

**Lemma 1.** *The  $\mathcal{Y}$ -space marginal vector field that generates the probability path  $p_t^Y$  takes the form*

$$\mathbf{u}_t^Y(\mathbf{y}) = \mathbf{A} \mathbf{v}_{\theta^*}^X(\mathbf{y}, t) - (a_t a'_t \sigma_0^2 + b_t b'_t \sigma_1^2) \nabla_{\mathbf{y}} \log p_t^Y(\mathbf{y}). \quad (17)$$

Meanwhile, the relationship between the  $\mathcal{Y}$ -space marginal vector field  $\mathbf{u}_t^Y(\mathbf{y})$  and the score function  $\nabla_{\mathbf{y}} \log p_t^Y(\mathbf{y})$  is established by the following lemma.

**Lemma 2.** *Note that  $p_1^Y(\mathbf{y}) = \int p_1(\mathbf{y} | \mathbf{x}) p_1^X(\mathbf{x}) d\mathbf{x} = \mathcal{CN}(\mathbf{y} | \mathbf{0}, 2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I}_{Cd})$ . Then,*

$$\mathbf{u}_t^Y(\mathbf{y}) = \frac{a'_t}{a_t} \mathbf{y} - b_t \left( b'_t - \frac{a'_t}{a_t} b_t \right) (2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I}_{Cd}) \nabla_{\mathbf{y}} \log p_t^Y(\mathbf{y}). \quad (18)$$

The proof is done by writing  $\mathbf{u}_t^Y(\mathbf{y})$  and  $\nabla_{\mathbf{y}} \log p_t^Y(\mathbf{y})$  in terms of the conditional expectation  $\mathbb{E}_{q_t(\mathbf{y}_0 | \mathbf{y})}[\mathbf{y}_1]$ , which can be found in Appendix A. Combining Lemma 1 and Lemma 2 by canceling out the score function gives the following proposition that relates the  $\mathcal{Y}$ -space marginal vector field to the optimal solution to PCFM.

**Proposition 3 (Vector fields under projection).** *For  $a_t = 1 - t$  and  $b_t = t$ , the  $\mathcal{Y}$ -space marginal vector field  $\mathbf{u}_t^Y(\mathbf{y})$  can be expressed by  $\mathbf{v}_{\theta^*}^X(\mathbf{y}, t)$  as*

$$\mathbf{u}_t^Y(\mathbf{y}) = \mathbf{A} \mathbf{v}_{\theta^*}^X(\mathbf{y}, t) - \frac{c_t}{1-t} [(c_t + \sigma_1^2) \mathbf{I}_{Cd} + 2\mathbf{A}\mathbf{A}^*]^{-1} [(1-t) \mathbf{A} \mathbf{v}_{\theta^*}^X(\mathbf{y}, t) + \mathbf{y}], \quad (19)$$

where  $c_t \triangleq (1-t) \left( \frac{1-t}{t} \sigma_0^2 - \sigma_1^2 \right)$ . In addition, left-multiplying both sides with  $\mathbf{A}^*$  gives the more computationally friendly formula when  $Cd > D$ :

$$\mathbf{A}^* \mathbf{u}_t^Y(\mathbf{y}) = \mathbf{A}^* \mathbf{A} \mathbf{v}_{\theta^*}^X(\mathbf{y}, t) - \frac{c_t}{1-t} [(c_t + \sigma_1^2) \mathbf{I}_D + 2\mathbf{A}^* \mathbf{A}]^{-1} \mathbf{A}^* [(1-t) \mathbf{A} \mathbf{v}_{\theta^*}^X(\mathbf{y}, t) + \mathbf{y}]. \quad (20)$$

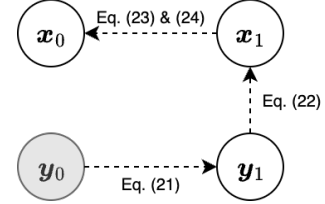


Figure 2: Inference steps of the proposed reconstruction algorithm.

**Algorithm 1:** Reconstruction via Dual-Space Cyclic Integration with PCFM

---

**Input:** k-space measurement  $\mathbf{y}_0$ , pretrained optimal solution to PCFM  $\mathbf{v}_{\theta^*}^X(\cdot, t)$ , number of time steps  $T$ .  
**Output:** Reconstructed image  $\mathbf{x}_0$  of  $\mathbf{y}_0$ .

---

```

1 for  $t = 0, \dots, (T-1)/T$  do
2    $\mathbf{y}_{t+1/T} \leftarrow \mathbf{y}_t + \frac{1}{T} \mathbf{u}_t^Y(\mathbf{y})$  ▷ Forward integration using Proposition 3
3   Sample  $\mathbf{x}_1 \sim p_1(\mathbf{x} \mid \mathbf{y})$ . ▷ Posterior sampling with Eq. (22)
4   for  $t \in \{(T-1)/T, \dots, 0\}$  do
5      $\tilde{\mathbf{y}}_{t+1/T} \leftarrow a_{t+1/T} \mathbf{y}_0 + b_{t+1/T} \mathbf{y}_1$ 
6      $\tilde{\mathbf{x}}_t \leftarrow \mathbf{x}_{t+1/T} - \frac{1}{T} \mathbf{v}_{\theta^*}^X(\tilde{\mathbf{y}}_{t+1/T}, t + 1/T)$  ▷ Backward integration
7      $\mathbf{x}_t \leftarrow \tilde{\mathbf{x}}_t - \mathbf{A}^+(\mathbf{A}\tilde{\mathbf{x}}_t - \tilde{\mathbf{y}}_t)$  ▷ Data consistency step
8 return  $\mathbf{x}_0$ .
```

---

This proposition asserts that, using the optimal solution of the PCFM, an associated marginal vector field can be derived, which facilitates the generation of the probability flow within the measurement space. Consequently, to sample from  $p(\mathbf{y}_1 \mid \mathbf{y}_0)$ , the  $\mathcal{Y}$ -space flow can be forward integrated using  $\mathbf{u}_t^Y$  to produce the subsequent  $\mathbf{y}_1$  from the observed initial  $\mathbf{y}_0$ , i.e.,

$$d\mathbf{y}_t = \mathbf{u}_t^Y(\mathbf{y}_t) dt. \quad (21)$$

**Posterior sampling.** The posterior  $p(\mathbf{x}_1 \mid \mathbf{y}_1)$  follows a closed-form Gaussian distribution

$$\begin{aligned} p(\mathbf{x}_1 \mid \mathbf{y}_1) &= p_1(\mathbf{x} \mid \mathbf{y}) \propto p_1^X(\mathbf{x}) p_1(\mathbf{y} \mid \mathbf{x}) \\ &= \mathcal{CN}\left(\mathbf{x} \mid \left(\frac{\sigma_1^2}{2} \mathbf{I}_D + \mathbf{A}^* \mathbf{A}\right)^{-1} \mathbf{A}^* \mathbf{y}, \sigma_1^2 \left(\frac{\sigma_1^2}{2} \mathbf{I}_D + \mathbf{A}^* \mathbf{A}\right)^{-1}\right), \end{aligned} \quad (22)$$

which can be easily sampled by linear transformation of a standard Gaussian vector.

**Backward integration and measurement consistency.** The distribution  $p(\mathbf{x}_0 \mid \mathbf{x}_1)$  is a delta distribution given  $\mathbf{x}_1$  due to the  $\mathcal{X}$ -space flow, which we can approximate using the backward ODE integration with the PCFM optimal solution, i.e.,

$$d\tilde{\mathbf{x}}_t = \mathbf{v}_{\theta^*}^X(\tilde{\mathbf{y}}_t, t) dt, \quad (23)$$

where  $\tilde{\mathbf{x}}_1 \triangleq \mathbf{x}_1$  and  $\tilde{\mathbf{y}}_t \triangleq a_t \mathbf{y}_0 + b_t \mathbf{y}_1$ . In addition, to enforce measurement consistency with  $\tilde{\mathbf{y}}_t$  as in many inverse problem solving algorithms (Daras et al., 2024a), we can employ the range-null decomposition (Wang et al., 2023) after each step of the backward integration if the observed measurement  $\mathbf{y}_0$  is clean, i.e.,

$$\mathbf{x}_t \leftarrow \tilde{\mathbf{x}}_t - \mathbf{A}^+(\mathbf{A}\tilde{\mathbf{x}}_t - \tilde{\mathbf{y}}_t), \quad (24)$$

where the pseudoinverse operator  $\mathbf{A}^+$  can be approximated by the CG method. However, if the measurement  $\mathbf{y}_0$  is noisy, this step will also inject noise into the reconstruction. To address this, we find that the Plug-and-Play (PnP) framework can perform better on noisy measurement (Combettes & Wajs, 2005; Venkatakrishnan et al., 2013; Martin et al., 2025), which we introduce in Appendix B. Ablation study on the backward sampling strategies is presented in Appendix E.2.

**Reconstruction algorithm.** The overall discrete-time algorithm is outlined in Alg. 1 for scenarios involving noiseless measurements and in Alg. 2 within Appendix B for noisy measurements.

## 4 RELATED WORK

**Diffusion model & flow matching for inverse problems.** A comprehensive review on diffusion models for inverse problems is provided by (Daras et al., 2024a) and (Chung et al., 2025). For example, diffusion posterior sampling (DPS) is proposed to sample from the posterior distribution based on score matching (Chung et al., 2023) which however requires the number of function evaluations (NFEs)=1000. (Wang et al., 2023) proposed range-null decomposition based on DDIM sampling

(Song et al., 2021), which only needs NFEs=100. Recently, flow matching has also been explored to solve inverse problems (Pokle et al., 2024; Zhang et al., 2024; Qin et al., 2025; Martin et al., 2025; Yan et al., 2025). For example, (Pokle et al., 2024) proposed to combine the prior score function based on flow matching and the likelihood score based on IIGDM (Song et al., 2023). (Martin et al., 2025) proposed to integrate the PnP framework with flow matching. Nevertheless, existing diffusion models and flow matching require large amounts of ground-truth data to learn their prior distribution, which are impossible or expensive to acquire for MRI in real scenarios, for example, in dynamic or low-field MRI (Lustig et al., 2007; Marques et al., 2019).

**Self-supervised & unsupervised MRI reconstruction.** We explicitly distinguish between self-supervised methods based on additional subsampling of the available k-space and unsupervised approaches that use all the observed undersampled k-space measurement, while both are referred to as self-supervised in a recent benchmark study (Wang et al., 2025). Multiple recent studies follow the prior learning paradigm and have explored methodologies to infer the prior distribution of ground-truth data from only corrupted measurements based on diffusion models or flow matching. The ambient diffusion family (Daras et al., 2023; Aali et al., 2025; Daras et al., 2024b; 2025) proposed optimizing the denoising diffusion objective through a self-supervised approach by introducing further corruption to the measurements. Our work is more related to (Kawar et al., 2024) and (Luo et al., 2025) where they proposed adapting the Ensemble SURE frameworks (Aggarwal et al., 2022) to the diffusion model and flow matching objectives in an unsupervised fashion. They focused solely on a single-coil MRI model with a feasible weighted projection operator. In contrast, we introduce a rigorous GSURE-based projected CFM formulation that facilitates optimization in parallel MRI.

## 5 NUMERICAL EXPERIMENTS

### 5.1 EXPERIMENTAL SETUPS

**Datasets and preprocessing.** Experiments were conducted on the NYU fastMRI (Knoll et al., 2020; Zbontar et al., 2018) and the CMRxRecon challenge (2023) (Wang et al., 2024; Lyu et al., 2025) datasets to evaluate the performance of the model for accelerated multi-coil MRI reconstruction. A selection of 11094/1584/3172 T2-weighted brain MRI slices and 5451/779/1557 cardiac T1/T2 quantitative mapping slices was randomly sampled for training/validation/test, respectively. Ground-truth images were obtained by the SENSE reconstruction from fully sampled k-space (Pruessmann et al., 1999), i.e.,  $\mathbf{x}_0 \triangleq (\sum_c \mathbf{S}_c^* \mathbf{S}_c)^{-1} \sum_c \mathbf{S}_c^* \mathbf{F}^* \hat{\mathbf{y}}_{0,c}$ , where  $\hat{\mathbf{y}}_0$  is the fully sampled k-space measurement, and the coil sensitivity maps were estimated by ESPIRiT (Uecker et al., 2014). We retrospectively simulated a random Cartesian (1D) undersampling mask for each image, where every mask contains fully sampled low-frequency k-space lines. The other lines were randomly uniformly sampled according to the acceleration factor. The zero-filled adjoint transform  $\mathbf{A}^* \mathbf{y}_0$  is the undersampled image before reconstruction.

**Implementation details.** The conditional OT path (Lipman et al., 2023) is adopted for the proposed framework, where  $a_t = t$  and  $b_t = 1 - t$ . The noise level of the original k-space is assumed to be  $\sigma_0 = 10^{-3}$ , which can be considered as noiseless. We set  $\sigma_1 = 0$  for simplicity. We use ADM U-Net (Dhariwal & Nichol, 2021) as the network architecture for velocity field prediction, where each intermediate convolutional block is followed by adaptive group normalization whose parameters are conditioned on the time points. Multi-head attention and dropout layers are applied at the lowest three resolutions of the network. We train the network from scratch on the training data by the AdamW optimizer (Loshchilov & Hutter, 2019) for 100K steps, with a learning rate of  $10^{-4}$  and a weight decay coefficient of 0.1. Exponential moving average of the network parameters was performed every 100 training steps with a rate of 0.99. In the inference phase, we set the number of time steps as  $T = 10$  and the number of CG steps (Appendix C.1) for solving  $\mathbf{A}^+$  as 30.

### 5.2 BENCHMARK STUDY ON PUBLIC DATASETS

We benchmark our method against three types of baseline approaches on the fastMRI brain and CMRxRecon datasets. (A) Supervised methods that require fully sampled MRIs during training: **MoDL** (Aggarwal et al., 2018), **DDNM<sup>+</sup>** (Wang et al., 2023), **OT-ODE** (Pokle et al., 2024), and **PnP-Flow**, (Martin et al., 2025). (B) Self-supervised methods that learn to reconstruct by additional subsampling of the available k-space measurements: **SSDU** (Yaman et al., 2020), **Weighted SSDU** (Millard & Chiew, 2023) and **Robust SSDU** (Millard & Chiew, 2024). (C) Unsupervised methods

Table 1: Qualitative results of  $4\times$  and  $8\times$  accelerated multi-coil MRI reconstruction using various reconstruction methods on the fastMRI brain data with the random uniform undersampling pattern. Best results within each supervision category are highlighted in **bold**. The difference in metrics is statistically significant between our method and the others by the two-sided paired  $t$ -test ( $p < 0.05$ ).

| Training Supervision | Reconstruction Method | SSIM $\uparrow$                     |                                     | PSNR $\uparrow$                    |                                    | NFEs $\downarrow$ |
|----------------------|-----------------------|-------------------------------------|-------------------------------------|------------------------------------|------------------------------------|-------------------|
|                      |                       | $4\times$                           | $8\times$                           | $4\times$                          | $8\times$                          |                   |
| None                 | Zero-Filled           | $0.815 \pm 0.087$                   | $0.730 \pm 0.113$                   | $28.28 \pm 3.80$                   | $24.44 \pm 3.98$                   | N/A               |
| Supervised           | MoDL                  | <b><math>0.970 \pm 0.036</math></b> | $0.916 \pm 0.044$                   | $39.71 \pm 2.93$                   | $32.32 \pm 3.07$                   | 1                 |
|                      | DDNM <sup>+</sup>     | $0.938 \pm 0.041$                   | <b><math>0.920 \pm 0.040</math></b> | <b><math>40.00 \pm 2.78</math></b> | <b><math>35.24 \pm 2.85</math></b> | 100               |
|                      | OT-ODE                | $0.907 \pm 0.054$                   | $0.852 \pm 0.060$                   | $33.95 \pm 2.32$                   | $28.86 \pm 2.94$                   | 100               |
|                      | PnP-Flow              | $0.951 \pm 0.044$                   | $0.913 \pm 0.043$                   | $37.92 \pm 2.44$                   | $32.85 \pm 2.80$                   | 100               |
| Self-supervised      | SSDU                  | $0.831 \pm 0.076$                   | $0.792 \pm 0.094$                   | $28.13 \pm 2.44$                   | $26.61 \pm 3.68$                   | 1                 |
|                      | Weighted SSDU         | $0.939 \pm 0.049$                   | <b><math>0.870 \pm 0.055</math></b> | $36.09 \pm 2.64$                   | <b><math>29.72 \pm 2.90</math></b> | 1                 |
|                      | Robust SSDU           | <b><math>0.941 \pm 0.047</math></b> | $0.856 \pm 0.063$                   | <b><math>36.23 \pm 2.58</math></b> | $29.72 \pm 2.90$                   | 1                 |
| Unsupervised         | REI                   | $0.683 \pm 0.120$                   | $0.715 \pm 0.118$                   | $21.59 \pm 2.90$                   | $21.62 \pm 2.76$                   | 1                 |
|                      | MOI                   | $0.869 \pm 0.065$                   | $0.747 \pm 0.105$                   | $30.97 \pm 2.74$                   | $25.16 \pm 3.94$                   | 1                 |
|                      | ENSURE                | $0.899 \pm 0.065$                   | $0.800 \pm 0.104$                   | $32.75 \pm 4.20$                   | $27.12 \pm 4.33$                   | 1                 |
|                      | GTF <sup>2</sup> M    | $0.916 \pm 0.055$                   | $0.852 \pm 0.057$                   | $33.99 \pm 2.33$                   | $28.40 \pm 3.05$                   | 20                |
|                      | PCFM (Ours)           | <b><math>0.983 \pm 0.032</math></b> | <b><math>0.948 \pm 0.034</math></b> | <b><math>42.19 \pm 3.71</math></b> | <b><math>35.08 \pm 3.35</math></b> | 20                |

Table 2: Qualitative results of  $4\times$  and  $8\times$  accelerated multi-coil MRI reconstruction using various reconstruction methods on the CMRxRecon 2023 cardiac T1/T2 quantitative mapping data with the random uniform undersampling pattern. Best results within each supervision category are highlighted in **bold**. The difference in metrics is statistically significant between our method and the others by the two-sided paired  $t$ -test ( $p < 0.05$ ).

| Training Supervision | Reconstruction Method | SSIM $\uparrow$                     |                                     | PSNR $\uparrow$                    |                                    | NFEs $\downarrow$ |
|----------------------|-----------------------|-------------------------------------|-------------------------------------|------------------------------------|------------------------------------|-------------------|
|                      |                       | $4\times$                           | $8\times$                           | $4\times$                          | $8\times$                          |                   |
| None                 | Zero-Filled           | $0.769 \pm 0.057$                   | $0.747 \pm 0.056$                   | $27.01 \pm 1.97$                   | $26.11 \pm 1.90$                   | N/A               |
| Supervised           | MoDL                  | <b><math>0.979 \pm 0.011</math></b> | $0.943 \pm 0.025$                   | $41.79 \pm 3.38$                   | $36.36 \pm 3.14$                   | 1                 |
|                      | DDNM <sup>+</sup>     | $0.977 \pm 0.011$                   | <b><math>0.953 \pm 0.022</math></b> | <b><math>44.93 \pm 3.35</math></b> | <b><math>39.60 \pm 3.31</math></b> | 100               |
|                      | OT-ODE                | $0.922 \pm 0.036$                   | $0.870 \pm 0.052$                   | $35.65 \pm 3.23$                   | $32.08 \pm 3.00$                   | 100               |
|                      | PnP-Flow              | $0.963 \pm 0.021$                   | $0.924 \pm 0.038$                   | $39.77 \pm 3.57$                   | $35.13 \pm 3.46$                   | 100               |
| Self-supervised      | SSDU                  | $0.888 \pm 0.035$                   | $0.811 \pm 0.051$                   | $32.64 \pm 2.59$                   | $29.23 \pm 2.37$                   | 1                 |
|                      | Weighted SSDU         | $0.872 \pm 0.050$                   | $0.831 \pm 0.050$                   | $32.11 \pm 2.88$                   | $29.82 \pm 2.58$                   | 1                 |
|                      | Robust SSDU           | <b><math>0.912 \pm 0.031</math></b> | <b><math>0.861 \pm 0.044</math></b> | <b><math>33.93 \pm 2.59</math></b> | <b><math>31.10 \pm 2.66</math></b> | 1                 |
| Unsupervised         | REI                   | $0.736 \pm 0.056$                   | $0.716 \pm 0.060$                   | $23.65 \pm 2.12$                   | $26.62 \pm 1.96$                   | 1                 |
|                      | MOI                   | $0.971 \pm 0.015$                   | $0.874 \pm 0.040$                   | $40.13 \pm 3.26$                   | $31.47 \pm 2.62$                   | 1                 |
|                      | ENSURE                | $0.918 \pm 0.028$                   | $0.849 \pm 0.052$                   | $34.57 \pm 3.13$                   | $30.27 \pm 2.48$                   | 1                 |
|                      | GTF <sup>2</sup> M    | $0.918 \pm 0.028$                   | $0.881 \pm 0.038$                   | $33.67 \pm 2.43$                   | $31.33 \pm 2.42$                   | 20                |
|                      | PCFM (Ours)           | <b><math>0.994 \pm 0.008</math></b> | <b><math>0.974 \pm 0.016</math></b> | <b><math>52.39 \pm 9.81</math></b> | <b><math>40.50 \pm 3.84</math></b> | 20                |

that learn to reconstruct using all available k-space measurements: **REI** (Chen et al., 2022), **MOI** (Tachella et al., 2022), **ENSURE** (Aggarwal et al., 2022), and **GTF<sup>2</sup>M** (Luo et al., 2025). Details of the baseline setups can be found in Appendix D.

Table 1 and Table 2 present quantitative results on the test data of multi-coil brain and cardiac MRI, respectively. The tables show that our method ranks first in terms of SSIM and PSNR on both datasets among all self-supervised and unsupervised methods, and outperforms all supervised approaches except for PSNR on fastMRI  $8\times$  data. Furthermore, our approach demonstrates better efficiency relative to baseline approaches based on supervised diffusion models or flow matching, even though it employs the same network structure and training strategy. This is achieved with a significant decrease in NFEs when compared to supervised diffusion models and flow matching-based techniques. Fig. 3 and Fig. S1 visualize reconstruction and error map on several multi-coil brain and cardiac MRI test samples, respectively. Our method produces error maps that either match or exceed those generated by supervised baseline approaches, particularly in areas around anatomical boundaries where high-frequency details are absent in the undersampled k-space. Additional ablation studies on forward/backward sampling strategies and inference time comparison are provided in Appendix E.

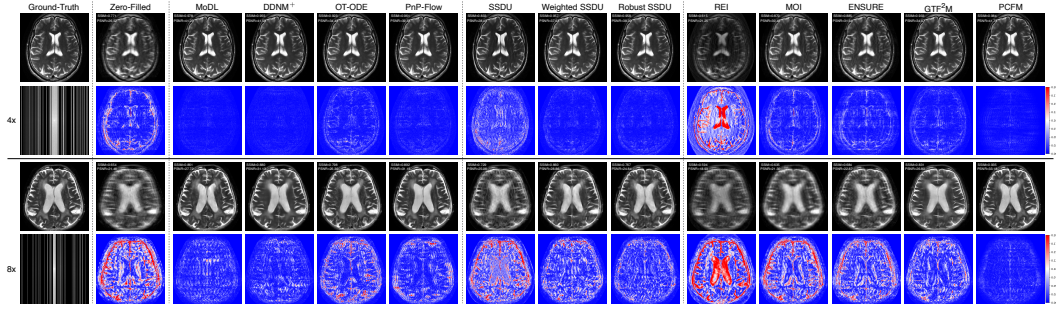


Figure 3: Visualization of reconstruction on two test samples of  $4\times$  and  $8\times$  accelerated multi-coil brain MRI from the compared methods. The k-space are presented in log-scale absolute values. The error maps are presented in values relative to the peak intensity in the ground-truth image.

Table 3: Qualitative results of  $4\times$  and  $8\times$  reconstruction of CMRxRecon 2023 cardiac T1/T2 quantitative mapping images trained and tested both on noisy data. The difference in metrics is statistically significant between our method and the others by the two-sided paired  $t$ -test ( $p < 0.05$ ).

| Training Noise level | Reconstruction Method | SSIM $\uparrow$                     |                                     | PSNR $\uparrow$                    |                                    | NFEs $\downarrow$ |
|----------------------|-----------------------|-------------------------------------|-------------------------------------|------------------------------------|------------------------------------|-------------------|
|                      |                       | $4\times$                           | $8\times$                           | $4\times$                          | $8\times$                          |                   |
| $\sigma_0 = 0.05$    | Zero-Filled           | $0.768 \pm 0.057$                   | $0.746 \pm 0.056$                   | $27.01 \pm 1.97$                   | $26.11 \pm 1.90$                   | N/A               |
|                      | Robust SSDU           | $0.912 \pm 0.031$                   | $0.855 \pm 0.045$                   | $34.06 \pm 2.69$                   | $30.81 \pm 2.61$                   | 10                |
|                      | ENSURE                | $0.861 \pm 0.050$                   | $0.800 \pm 0.058$                   | $33.25 \pm 2.75$                   | $29.80 \pm 2.39$                   | 1                 |
|                      | PnP-PCFM (Ours)       | <b><math>0.928 \pm 0.027</math></b> | <b><math>0.892 \pm 0.038</math></b> | <b><math>35.58 \pm 2.75</math></b> | <b><math>33.14 \pm 2.82</math></b> | 20                |
| $\sigma_0 = 0.1$     | Zero-Filled           | $0.763 \pm 0.058$                   | $0.744 \pm 0.057$                   | $26.98 \pm 1.97$                   | $26.10 \pm 1.90$                   | N/A               |
|                      | Robust SSDU           | $0.906 \pm 0.033$                   | $0.855 \pm 0.045$                   | $33.80 \pm 2.68$                   | $30.91 \pm 2.64$                   | 10                |
|                      | ENSURE                | $0.784 \pm 0.063$                   | $0.769 \pm 0.062$                   | $31.39 \pm 2.54$                   | $28.85 \pm 2.30$                   | 1                 |
|                      | PnP-PCFM (Ours)       | <b><math>0.921 \pm 0.029</math></b> | <b><math>0.889 \pm 0.039</math></b> | <b><math>35.62 \pm 2.85</math></b> | <b><math>33.04 \pm 2.82</math></b> | 20                |
| $\sigma_0 = 0.2$     | Zero-Filled           | $0.744 \pm 0.062$                   | $0.732 \pm 0.060$                   | $26.89 \pm 1.96$                   | $26.05 \pm 1.91$                   | N/A               |
|                      | Robust SSDU           | $0.878 \pm 0.039$                   | $0.822 \pm 0.049$                   | $32.47 \pm 2.52$                   | $29.78 \pm 2.39$                   | 10                |
|                      | ENSURE                | $0.563 \pm 0.094$                   | $0.542 \pm 0.090$                   | $25.88 \pm 2.55$                   | $25.29 \pm 2.39$                   | 1                 |
|                      | PnP-PCFM (Ours)       | <b><math>0.885 \pm 0.038</math></b> | <b><math>0.868 \pm 0.044</math></b> | <b><math>34.34 \pm 2.76</math></b> | <b><math>32.52 \pm 2.77</math></b> | 20                |

### 5.3 NOISY DATA

We proceed to assess our method on data with noise by introducing additive Gaussian noise with  $\sigma_0 = 0.05, 0.1, 0.2$  to the training set based on the original k-space measurements. The model’s performance is evaluated on both noiseless and noisy test datasets. For evaluation, Alg. 1 is used if the data set is noiseless, otherwise Alg. 2 is used. Table S4 displays the quantitative results regarding the noiseless test data derived from the CMRxRecon 2023 dataset. A notable observation is that the model retains its performance as if the training data were free of noise, suggesting that PCFM effectively learns the prior distribution of fully sampled MRI even when trained on noisy and undersampled inputs. Table 3 presents the quantitative results for noisy test data, comparing PnP-PCFM against the baseline models Robust SSDU and ENSURE, which are designed to handle noisy inputs. The proposed approach exhibits superior performance compared to them. Fig. S2 provides a visualization of reconstruction and error maps for two test samples involving  $4\times$  and  $8\times$  accelerated multi-coil cardiac MRI with varying levels of additive Gaussian noise.

## 6 CONCLUSION

In this study, we present Projected Conditional Flow Matching (PCFM), a new framework that utilizes the generalized Stein’s unbiased risk estimator to learn the prior distribution of fully sampled parallel MRI using solely undersampled k-space data. From the optimal PCFM solution, we derived a marginal vector field within the associated measurement space, facilitating the creation of a dual-space cyclic integration method for MRI reconstruction. Experimental assessments on two parallel MRI datasets demonstrate that PCFM attains leading performance relative to prior baselines on both noiseless and noisy data.

## DECLARATIONS

**LLM Usage:** We note that large language models (LLMs) were used solely to improve language clarity; all scientific content and ideas were developed by the authors.

**Ethics Statement:** This research complies with the ICLR Code of Ethics. The study uses multiple datasets: publicly available fastMRI and CMRxRecon 2023. All datasets are de-identified and contain no personally identifiable information. No direct human subject recruitment was involved, and the work does not raise foreseeable risks to participants.

**Reproducibility Statement:** We provide detailed descriptions of datasets, preprocessing steps, model architectures, training procedures, and evaluation protocols in Section 3 and Section 5 of the manuscript and Appendix B, Appendix C, and Appendix D of the supplementary material to ensure reproducibility of our model and compared approaches. Source code and model weights will be made publicly available at <https://github.com/anonymous> upon acceptance.

## REFERENCES

- Asad Aali, Giannis Daras, Brett Levac, Sidharth Kumar, Alexandros G Dimakis, and Jonathan I Tamir. Ambient diffusion posterior sampling: Solving inverse problems with diffusion models trained on corrupted data. In *International Conference on Learning Representations*, 2025.
- Hemant K Aggarwal, Merry P Mani, and Mathews Jacob. MoDL: Model-based deep learning architecture for inverse problems. *IEEE Transactions on Medical Imaging*, 38(2):394–405, 2018.
- Hemant Kumar Aggarwal, Aniket Pramanik, Maneesh John, and Mathews Jacob. ENSURE: A general approach for unsupervised training of deep image reconstruction algorithms. *IEEE Transactions on Medical Imaging*, 42(4):1133–1144, 2022.
- Amir Beck. *First-order methods in optimization*. SIAM, 2017.
- Stephen Boyd, Neal Parikh, Eric Chu, Borja Peleato, Jonathan Eckstein, et al. Distributed optimization and statistical learning via the alternating direction method of multipliers. *Foundations and Trends® in Machine Learning*, 3(1):1–122, 2011.
- Dongdong Chen, Julián Tachella, and Mike E Davies. Equivariant imaging: Learning beyond the range space. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pp. 4379–4388, 2021.
- Dongdong Chen, Julián Tachella, and Mike E Davies. Robust equivariant imaging: a fully unsupervised framework for learning to image from noisy and partial measurements. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 5647–5656, 2022.
- Ricky TQ Chen, Yulia Rubanova, Jesse Bettencourt, and David K Duvenaud. Neural ordinary differential equations. *Advances in Neural Information Processing systems*, 31, 2018.
- Hyungjin Chung and Jong Chul Ye. Score-based diffusion models for accelerated mri. *Medical Image Analysis*, 80:102479, 2022.
- Hyungjin Chung, Jeongsol Kim, Michael Thompson Mccann, Marc Louis Klasky, and Jong Chul Ye. Diffusion posterior sampling for general noisy inverse problems. In *International Conference on Learning Representations*, 2023.
- Hyungjin Chung, Jeongsol Kim, and Jong Chul Ye. Diffusion models for inverse problems. *arXiv preprint arXiv:2508.01975*, 2025.
- Patrick L Combettes and Valérie R Wajs. Signal recovery by proximal forward-backward splitting. *Multiscale Modeling & Simulation*, 4(4):1168–1200, 2005.
- Giannis Daras, Kulin Shah, Yuval Dagan, Aravind Gollakota, Alex Dimakis, and Adam Klivans. Ambient diffusion: Learning clean distributions from corrupted data. *Advances in Neural Information Processing Systems*, 36:288–313, 2023.

- Giannis Daras, Hyungjin Chung, Chieh-Hsin Lai, Yuki Mitsufuji, Jong Chul Ye, Peyman Milanfar, Alexandros G Dimakis, and Mauricio Delbracio. A survey on diffusion models for inverse problems. *arXiv preprint arXiv:2410.00083*, 2024a.
- Giannis Daras, Alexandros G Dimakis, and Constantinos Daskalakis. Consistent diffusion meets tweedie: Training exact ambient diffusion models with noisy data. In *International Conference on Machine Learning*, 2024b.
- Giannis Daras, Adrian Rodriguez-Munoz, Adam Klivans, Antonio Torralba, and Constantinos Daskalakis. Ambient diffusion omni: Training good models with bad data. *arXiv preprint arXiv:2506.10038*, 2025.
- Prafulla Dhariwal and Alexander Nichol. Diffusion models beat gans on image synthesis. *Advances in Neural Information Processing Systems*, 34:8780–8794, 2021.
- Yonina C Eldar. Generalized sure for exponential families: Applications to regularization. *IEEE Transactions on Signal Processing*, 57(2):471–481, 2008.
- Zhengan Fang, Sam Buchanan, and Jeremias Sulam. What’s in a prior? learned proximal networks for inverse problems. In *International Conference on Learning Representations*, 2024.
- Donald Geman and Chengda Yang. Nonlinear image recovery with half-quadratic regularization. *IEEE Transactions on Image Processing*, 4(7):932–946, 1995.
- Mark A Griswold, Peter M Jakob, Robin M Heidemann, Mathias Nittka, Vladimir Jellus, Jianmin Wang, Berthold Kiefer, and Axel Haase. Generalized autocalibrating partially parallel acquisitions (grappa). *Magnetic Resonance in Medicine*, 47(6):1202–1210, 2002.
- Kerstin Hammernik, Teresa Klatzer, Erich Kobler, Michael P Recht, Daniel K Sodickson, Thomas Pock, and Florian Knoll. Learning a variational network for reconstruction of accelerated mri data. *Magnetic Resonance in Medicine*, 79(6):3055–3071, 2018.
- Jonathan Ho, Ajay Jain, and Pieter Abbeel. Denoising diffusion probabilistic models. *Advances in Neural Information Processing Systems*, 33:6840–6851, 2020.
- Bahjat Kavar, Noam Elata, Tomer Michaeli, and Michael Elad. Gsure-based diffusion model training with corrupted data. *Transactions on Machine Learning Research*, 2024.
- Florian Knoll, Jure Zbontar, Anuroop Sriram, Matthew J Muckley, Mary Bruno, Aaron Defazio, Marc Parente, Krzysztof J Geras, Joe Katsnelson, Hersh Chandarana, et al. fastmri: A publicly available raw k-space and dicom dataset of knee images for accelerated mr image reconstruction using machine learning. *Radiology: Artificial Intelligence*, 2(1):e190007, 2020.
- Yaron Lipman, Ricky TQ Chen, Heli Ben-Hamu, Maximilian Nickel, and Matthew Le. Flow matching for generative modeling. In *International Conference on Learning Representations*, 2023.
- Xingchao Liu, Chengyue Gong, and Qiang Liu. Flow straight and fast: Learning to generate and transfer data with rectified flow. In *International Conference on Learning Representations*, 2023.
- Ilya Loshchilov and Frank Hutter. Decoupled weight decay regularization. In *International Conference on Learning Representations*, 2019.
- Xinzhe Luo, Yingzhen Li, and Chen Qin. Unsupervised accelerated mri reconstruction via ground-truth-free flow matching. In *International Conference on Information Processing in Medical Imaging*, pp. 127–141. Springer, 2025.
- Michael Lustig, David Donoho, and John M Pauly. Sparse mri: The application of compressed sensing for rapid mr imaging. *Magnetic Resonance in Medicine*, 58(6):1182–1195, 2007.
- Michael Lustig, David L Donoho, Juan M Santos, and John M Pauly. Compressed sensing mri. *IEEE Signal Processing Magazine*, 25(2):72–82, 2008.
- Jun Lyu, Chen Qin, Shuo Wang, Fanwen Wang, Yan Li, Zi Wang, Kunyuan Guo, Cheng Ouyang, Michael Tänzler, Meng Liu, et al. The state-of-the-art in cardiac mri reconstruction: Results of the cmrxrecon challenge in miccai 2023. *Medical Image Analysis*, 101:103485, 2025.

- Morteza Mardani, Enhao Gong, Joseph Y Cheng, Shreyas S Vasanawala, Greg Zaharchuk, Lei Xing, and John M Pauly. Deep generative adversarial neural networks for compressive sensing mri. *IEEE Transactions on Medical Imaging*, 38(1):167–179, 2018.
- José P Marques, Frank FJ Simonis, and Andrew G Webb. Low-field mri: an mr physics perspective. *Journal of Magnetic Resonance Imaging*, 49(6):1528–1542, 2019.
- Ségolène Martin, Anne Gagneux, Paul Hagemann, and Gabriele Steidl. Pnp-flow: Plug-and-play image restoration with flow matching. In *International Conference on Learning Representations*, 2025.
- Charles Millard and Mark Chiew. A theoretical framework for self-supervised mr image reconstruction using sub-sampling via variable density noisier2noise. *IEEE Transactions on Computational Imaging*, 9:707–720, 2023.
- Charles Millard and Mark Chiew. Clean self-supervised mri reconstruction from noisy, sub-sampled training data with robust ssdu. *Bioengineering*, 11(12):1305, 2024.
- Alexander Quinn Nichol and Prafulla Dhariwal. Improved denoising diffusion probabilistic models. In *International Conference on Machine Learning*, pp. 8162–8171. PMLR, 2021.
- Ashwini Pogle, Matthew J. Muckley, Ricky T. Q. Chen, and Brian Karrer. Training-free linear image inverses via flows. *Transactions on Machine Learning Research*, 2024. ISSN 2835-8856.
- Klaas P Pruessmann, Markus Weiger, Markus B Scheidegger, and Peter Boesiger. Sense: sensitivity encoding for fast mri. *Magnetic Resonance in Medicine*, 42(5):952–962, 1999.
- Haina Qin, Wenyang Luo, Libin Wang, Dandan Zheng, Jingdong Chen, Ming Yang, Bing Li, and Weiming Hu. Reversing flow for image restoration. In *Proceedings of the Computer Vision and Pattern Recognition Conference*, pp. 7545–7558, 2025.
- Peter B Roemer, William A Edelstein, Cecil E Hayes, Steven P Souza, and Otward M Mueller. The NMR phased array. *Magnetic Resonance in Medicine*, 16(2):192–225, 1990.
- Daniel K Sodickson and Warren J Manning. Simultaneous acquisition of spatial harmonics (smash): fast imaging with radiofrequency coil arrays. *Magnetic Resonance in Medicine*, 38(4):591–603, 1997.
- Jiaming Song, Chenlin Meng, and Stefano Ermon. Denoising diffusion implicit models. In *International Conference on Learning Representations*, 2021.
- Jiaming Song, Arash Vahdat, Morteza Mardani, and Jan Kautz. Pseudoinverse-guided diffusion models for inverse problems. In *International Conference on Learning Representations*, 2023.
- Yang Song, Liyue Shen, Lei Xing, and Stefano Ermon. Solving inverse problems in medical imaging with score-based generative models. In *International Conference on Learning Representations*, 2022.
- Charles M Stein. Estimation of the mean of a multivariate normal distribution. *The Annals of Statistics*, pp. 1135–1151, 1981.
- Julián Tachella, Dongdong Chen, and Mike Davies. Unsupervised learning from incomplete measurements for inverse problems. *Advances in Neural Information Processing Systems*, 35:4983–4995, 2022.
- Julián Tachella, Dongdong Chen, and Mike Davies. Sensing theorems for unsupervised learning in linear inverse problems. *Journal of Machine Learning Research*, 24(39):1–45, 2023.
- Kerem C Tezcan, Christian F Baumgartner, Roger Luechinger, Klaas P Pruessmann, and Ender Konukoglu. Mr image reconstruction using deep density priors. *IEEE Transactions on Medical Imaging*, 38(7):1633–1642, 2018.
- Alexander Tong, Kilian Fatras, Nikolay Malkin, Guillaume Hugué, Yanlei Zhang, Jarrid Rector-Brooks, Guy Wolf, and Yoshua Bengio. Improving and generalizing flow-based generative models with minibatch optimal transport. *Transactions on Machine Learning Research*, 2024, 2023.

- Martin Uecker, Peng Lai, Mark J Murphy, Patrick Virtue, Michael Elad, John M Pauly, Shreyas S Vasanawala, and Michael Lustig. ESPIRiT—an eigenvalue approach to autocalibrating parallel mri: where sense meets grappa. *Magnetic Resonance in Medicine*, 71(3):990–1001, 2014.
- Dmitry Ulyanov, Andrea Vedaldi, and Victor Lempitsky. Deep image prior. In *IEEE conference on Computer Vision and Pattern Recognition*, pp. 9446–9454, 2018.
- Singanallur V Venkatakrishnan, Charles A Bouman, and Brendt Wohlberg. Plug-and-play priors for model based reconstruction. In *IEEE Global Conference on Signal and Information Processing*, pp. 945–948. IEEE, 2013.
- Cédric Villani et al. *Optimal transport: old and new*, volume 338. Springer, 2009.
- Andrew Wang, Steven McDonagh, and Mike Davies. Benchmarking self-supervised learning methods for accelerated mri reconstruction. *arXiv preprint arXiv:2502.14009*, 2025.
- Chengyan Wang, Jun Lyu, Shuo Wang, Chen Qin, Kunyuan Guo, Xinyu Zhang, Xiaotong Yu, Yan Li, Fanwen Wang, Jianhua Jin, et al. Cmrrecon: A publicly available k-space dataset and benchmark to advance deep learning for cardiac mri. *Scientific Data*, 11(1):687, 2024.
- Yinhui Wang, Jiwen Yu, and Jian Zhang. Zero-shot image restoration using denoising diffusion null-space model. In *International Conference on Learning Representations*, 2023.
- Burhaneddin Yaman, Seyed Amir Hossein Hosseini, Steen Moeller, Jutta Ellermann, Kâmil Uğurbil, and Mehmet Akçakaya. Self-supervised learning of physics-guided reconstruction neural networks without fully sampled reference data. *Magnetic Resonance in Medicine*, 84(6):3172–3191, 2020.
- Yici Yan, Yichi Zhang, Xiangming Meng, and Zhizhen Zhao. Fig: Flow with interpolant guidance for linear inverse problems. In *International Conference on Learning Representations*, 2025.
- Jure Zbontar, Florian Knoll, Anuroop Sriram, Tullie Murrell, Zhengnan Huang, Matthew J Muckley, Aaron Defazio, Ruben Stern, Patricia Johnson, Mary Bruno, et al. fastmri: An open dataset and benchmarks for accelerated mri. *arXiv preprint arXiv:1811.08839*, 2018.
- Yasi Zhang, Peiyu Yu, Yaxuan Zhu, Yingshan Chang, Feng Gao, Ying Nian Wu, and Oscar Leong. Flow priors for linear inverse problems via iterative corrupted trajectory matching. *Advances in Neural Information Processing Systems*, 37:57389–57417, 2024.

# Appendices

## A PROOFS

**Proposition 1 (Optimal solution to PCFM).** *The minimizer of the PCFM objective is given by*

$$\mathbf{v}_{\theta^*}^X(\mathbf{y}, t) = \mathbb{E}_{q_t(\mathbf{z}^X|\mathbf{y}), p_t^X(\mathbf{x}|\mathbf{z}^X)} [\mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X)] + \mathbf{w}, \quad (\text{S1})$$

where  $q_t(\mathbf{z}^X | \mathbf{y}) = \frac{q(\mathbf{z}^X)p_t^Y(\mathbf{y}|\mathbf{z}^X)}{p_t^Y(\mathbf{y})}$ , and  $\mathbf{w}$  is any vector in the null space of  $\mathbf{A}$ , i.e.,  $\mathbf{A}\mathbf{w} = \mathbf{0}$ . In particular, when  $\mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X) = a'_t\mathbf{x}_0 + b'_t\mathbf{x}_1$  that is independent of  $\mathbf{x}$ , we have

$$\mathbf{v}_{\theta^*}^X(\mathbf{y}, t) = \mathbb{E}_{q_t(\mathbf{z}^X|\mathbf{y})} [\mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X)] + \mathbf{w}. \quad (\text{S2})$$

*Proof.* Since  $q(\mathbf{z}^X)p_t^X(\mathbf{x} | \mathbf{z}^X)p_t^Y(\mathbf{y} | \mathbf{z}^X) = p_t^Y(\mathbf{y})q_t(\mathbf{z}^X | \mathbf{y})p_t^X(\mathbf{x} | \mathbf{z}^X)$  and by the law of total expectation, the PCFM objective can be written as

$$\begin{aligned} \mathcal{L}_{\text{PCFM}}(\theta) &\triangleq \mathbb{E}_{t, q(\mathbf{z}^X), p_t^X(\mathbf{x}|\mathbf{z}^X), p_t^Y(\mathbf{y}|\mathbf{z}^X)} \|\mathbf{P}[\mathbf{v}_{\theta}^X(\mathbf{y}, t) - \mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X)]\|_2^2 \\ &= \mathbb{E}_{t, p_t^Y(\mathbf{y})} \mathbb{E}_{q_t^Y(\mathbf{z}^X|\mathbf{y}), p_t^X(\mathbf{x}|\mathbf{z}^X)} \|\mathbf{P}[\mathbf{v}_{\theta}^X(\mathbf{y}, t) - \mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X)]\|_2^2. \end{aligned} \quad (\text{S3})$$

To minimize the total expectation, we can minimize the inner conditional expectation for each value of  $t$  and  $\mathbf{y}$  independently. Let  $\hat{\mathbf{u}}_t^X(\mathbf{y}) \triangleq \mathbb{E}_{q_t(\mathbf{z}^X|\mathbf{y}), p_t^X(\mathbf{x}|\mathbf{z}^X)} [\mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X)]$ . For fixed  $t$  and  $\mathbf{y}$ , the inner expectation can be transformed as

$$\begin{aligned} \mathbf{I}_{t, \mathbf{y}}(\theta) &\triangleq \mathbb{E}_{q_t^Y(\mathbf{z}^X|\mathbf{y}), p_t^X(\mathbf{x}|\mathbf{z}^X)} \|\mathbf{P}[\mathbf{v}_{\theta}^X(\mathbf{y}, t) - \mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X)]\|_2^2 \\ &= \mathbb{E}_{q_t^Y(\mathbf{z}^X|\mathbf{y}), p_t^X(\mathbf{x}|\mathbf{z}^X)} \|\mathbf{P}[\mathbf{v}_{\theta}^X(\mathbf{y}, t) - \hat{\mathbf{u}}_t^X(\mathbf{y}) + \hat{\mathbf{u}}_t^X(\mathbf{y}) - \mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X)]\|_2^2 \\ &= \mathbb{E}_{q_t^Y(\mathbf{z}^X|\mathbf{y}), p_t^X(\mathbf{x}|\mathbf{z}^X)} \left[ \|\mathbf{P}[\mathbf{v}_{\theta}^X(\mathbf{y}, t) - \hat{\mathbf{u}}_t^X(\mathbf{y})]\|_2^2 \right. \\ &\quad \left. + 2(\mathbf{v}_{\theta}^X(\mathbf{y}, t) - \hat{\mathbf{u}}_t^X(\mathbf{y}))^* \mathbf{P}(\hat{\mathbf{u}}_t^X(\mathbf{y}) - \mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X)) \right] + \text{const.}, \end{aligned} \quad (\text{S4})$$

where we note that

$$\begin{aligned} &\mathbb{E}_{q_t^Y(\mathbf{z}^X|\mathbf{y}), p_t^X(\mathbf{x}|\mathbf{z}^X)} (\mathbf{v}_{\theta}^X(\mathbf{y}, t) - \hat{\mathbf{u}}_t^X(\mathbf{y}))^* \mathbf{P}(\hat{\mathbf{u}}_t^X(\mathbf{y}) - \mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X)) \\ &= (\mathbf{v}_{\theta}^X(\mathbf{y}, t) - \hat{\mathbf{u}}_t^X(\mathbf{y}))^* \mathbf{P} \left[ \hat{\mathbf{u}}_t^X(\mathbf{y}) - \mathbb{E}_{q_t^Y(\mathbf{z}^X|\mathbf{y}), p_t^X(\mathbf{x}|\mathbf{z}^X)} \mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X) \right] = 0. \end{aligned} \quad (\text{S5})$$

Therefore,  $\mathbf{I}_{t, \mathbf{y}}(\theta) = \mathbb{E}_{q_t^Y(\mathbf{z}^X|\mathbf{y}), p_t^X(\mathbf{x}|\mathbf{z}^X)} \|\mathbf{P}[\mathbf{v}_{\theta}^X(\mathbf{y}, t) - \hat{\mathbf{u}}_t^X(\mathbf{y})]\|_2^2$ , which is minimized when

$$\mathbf{v}_{\theta}^X(\mathbf{y}, t) = \hat{\mathbf{u}}_t^X(\mathbf{y}) + \mathbf{w}, \quad (\text{S6})$$

where  $\mathbf{w}$  is in the null space of  $\mathbf{P}$ , i.e.,  $\mathbf{P}\mathbf{w} = \mathbf{0}$ , which is equivalent to  $\mathbf{A}\mathbf{w} = \mathbf{0}$ .  $\square$

**Proposition 2 (Unsupervised transformation of PCFM).** *Assuming deterministic conditional probability paths  $\mathbf{x} = a_t\mathbf{x}_0 + b_t\mathbf{x}_1$  and  $\mathbf{y} = a_t\mathbf{y}_0 + b_t\mathbf{y}_1$  with  $\mathbf{y}_0 = \mathbf{A}\mathbf{x}_0 + \mathbf{e}_0$  and  $\mathbf{y}_1 = \mathbf{A}\mathbf{x}_1 + \mathbf{e}_1$ , where  $\mathbf{e}_0 \sim \mathcal{CN}(\mathbf{0}, \sigma_0^2 \mathbf{I}_{Cd})$  and  $\mathbf{e}_1 \sim \mathcal{CN}(\mathbf{0}, \sigma_1^2 \mathbf{I}_{Cd})$ , then up to a constant the PCFM objective can be transformed to*

$$\mathbb{E}_{t, q(\mathbf{z}^Y), p_t^Y(\mathbf{y}|\mathbf{z}^Y)} \left[ \|\mathbf{P}[\mathbf{v}_{\theta}^X(\mathbf{A}^*\mathbf{y}, t) - \hat{\mathbf{u}}_{t, ML}^X]\|_2^2 + \frac{2a_t}{a'_t} [(a'_t\sigma_0)^2 + (b'_t\sigma_1)^2] \nabla_{\mathbf{A}^*\mathbf{y}} \cdot \mathbf{P}\mathbf{v}_{\theta}^X(\mathbf{A}^*\mathbf{y}, t) \right], \quad (\text{S7})$$

where  $q(\mathbf{z}^Y) = q(\mathbf{y}_0)q(\mathbf{y}_1) = q(\mathbf{y}_0)\mathbb{E}_{q(\mathbf{x}_1)} [p(\mathbf{y}_1 | \mathbf{x}_1)]$  is sampled by the MRI and Monte Carlo estimation,  $\mathbf{P} = \mathbf{A}^+ \mathbf{A}$  is the range-space projection, and

$$\hat{\mathbf{u}}_{t, ML}^X \triangleq (\mathbf{A}^* \mathbf{C}_t^{-1} \mathbf{A})^+ \mathbf{A}^* \mathbf{C}_t^{-1} \mathbf{u}_t^Y(\mathbf{y} | \mathbf{z}^Y) \quad (\text{S8})$$

with  $\mathbf{C}_t = [(a'_t\sigma_0)^2 + (b'_t\sigma_1)^2] \mathbf{I}_d$  is the maximum likelihood solution of the forward model in Eq. (10). Note that  $\mathbf{A}^+$  denotes the Moore-Penrose pseudoinverse of  $\mathbf{A}$ , which can be approximated by the conjugate gradient method (Appendix C.1).

*Proof.* The proof is inspired from (Eldar, 2008). Noting the deterministic mapping between the  $\mathcal{Y}$ -space conditional path and the conditional vector field

$$\mathbf{y} = \frac{a_t}{a'_t} \mathbf{u}_t^Y(\mathbf{y} | \mathbf{z}^Y) - b'_t \left( \frac{a_t}{a'_t} - \frac{b_t}{b'_t} \right) \mathbf{y}_1, \quad (\text{S9})$$

we can write  $\mathbf{v}_\theta^X(\mathbf{y}, t) = \mathbf{v}_\theta^X(\mathbf{u}_t^Y(\mathbf{y} | \mathbf{z}^Y), t)$  and the objective as

$$\mathcal{L}_{\text{PCFM}}(\theta) = \mathbb{E}_{t, q(\mathbf{z}^X), p_t^X(\mathbf{x} | \mathbf{z}^X), p_t^Y(\mathbf{y} | \mathbf{z}^X)} \left\| \mathbf{P}[\mathbf{v}_\theta^X(\mathbf{u}_t^Y(\mathbf{y} | \mathbf{z}^Y), t) - \mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X)] \right\|_2^2. \quad (\text{S10})$$

By the linear forward model over the dual-space conditional vector fields

$$\mathbf{u}_t^Y(\mathbf{y} | \mathbf{z}^Y) = \mathbf{A} \mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X) + a'_t \mathbf{e}_0 + b'_t \mathbf{e}_1, \quad (\text{S11})$$

we note that the sufficient statistic  $\boldsymbol{\mu}_t^X \triangleq \mathbf{A}^* \mathbf{C}_t^{-1} \mathbf{u}_t^Y$  follows a Gaussian distribution  $\mathcal{CN}(\mathbf{A}^* \mathbf{C}_t^{-1} \mathbf{A} \mathbf{u}_t^X, \mathbf{A}^* \mathbf{C}_t^{-1} \mathbf{A})$  with probability density function (pdf)

$$p(\boldsymbol{\mu}_t^X) = q(\boldsymbol{\mu}_t^X) \exp \left( \mathbf{u}_t^{X*} \boldsymbol{\mu}_t^X - g(\mathbf{u}_t^X) \right), \quad (\text{S12})$$

where

$$\begin{aligned} q(\boldsymbol{\mu}_t^X) &= K \cdot \exp \left( -\frac{1}{2} \boldsymbol{\mu}_t^{X*} (\mathbf{A}^* \mathbf{C}_t^{-1} \mathbf{A})^+ \boldsymbol{\mu}_t^X \right), \\ g(\mathbf{u}_t^X) &= \frac{1}{2} \mathbf{u}_t^{X*} \mathbf{A}^* \mathbf{C}_t^{-1} \mathbf{A} \mathbf{u}_t^X. \end{aligned} \quad (\text{S13})$$

Assuming the network's input to be  $\boldsymbol{\mu}_t^X$ , we can write

$$\mathcal{L}_{\text{PCFM}}(\theta) = \mathbb{E}_{t, q(\mathbf{z}^X), p_t^X(\mathbf{x} | \mathbf{z}^X), p_t^Y(\mathbf{y} | \mathbf{z}^X)} \left[ \mathbf{u}_t^{X*} \mathbf{P} \mathbf{u}_t^X + \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t)^* \mathbf{P} \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t) - 2 \mathbf{u}_t^{X*} \mathbf{P} \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t) \right], \quad (\text{S14})$$

and

$$\begin{aligned} \mathbb{E}_{p_t^Y(\mathbf{y} | \mathbf{z}^X)} \left[ \mathbf{u}_t^{X*} \mathbf{P} \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t) \right] &= \mathbb{E}_{p(\boldsymbol{\mu}_t^X)} \left[ \mathbf{u}_t^{X*} \mathbf{P} \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t) \right] \\ &= \int \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t)^* \mathbf{P} \mathbf{u}_t^X \cdot q(\boldsymbol{\mu}_t^X) \exp \left( \mathbf{u}_t^{X*} \boldsymbol{\mu}_t^X - g(\mathbf{u}_t^X) \right) d\boldsymbol{\mu}_t^X. \end{aligned} \quad (\text{S15})$$

Denote  $h(\boldsymbol{\mu}_t^X) \triangleq \exp \left( \mathbf{u}_t^{X*} \boldsymbol{\mu}_t^X - g(\mathbf{u}_t^X) \right)$ . Substituting  $\mathbf{u}_t^X h(\boldsymbol{\mu}_t^X) = \nabla_{\boldsymbol{\mu}_t^X} h(\boldsymbol{\mu}_t^X)$  and integrating by parts, we have

$$\begin{aligned} \mathbb{E}_{p(\boldsymbol{\mu}_t^X)} \left[ \mathbf{u}_t^{X*} \mathbf{P} \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t) \right] &= \int \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t)^* \mathbf{P} \mathbf{u}_t^X \cdot q(\boldsymbol{\mu}_t^X) \nabla_{\boldsymbol{\mu}_t^X} h(\boldsymbol{\mu}_t^X) d\boldsymbol{\mu}_t^X \\ &= - \int h(\boldsymbol{\mu}_t^X) \nabla_{\boldsymbol{\mu}_t^X} \cdot [\mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t)^* \mathbf{P} \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t)] d\boldsymbol{\mu}_t^X, \end{aligned} \quad (\text{S16})$$

where

$$\nabla_{\boldsymbol{\mu}_t^X} \cdot [\mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t)^* \mathbf{P} \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t)] = q(\boldsymbol{\mu}_t^X) \left[ \nabla_{\boldsymbol{\mu}_t^X} \cdot \mathbf{P} \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t) + \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t)^* \mathbf{P} \nabla_{\boldsymbol{\mu}_t^X} \ln q(\boldsymbol{\mu}_t^X) \right] \quad (\text{S17})$$

and  $\ln q(\boldsymbol{\mu}_t^X) = -(\mathbf{A}^* \mathbf{C}_t^{-1} \mathbf{A})^+ \boldsymbol{\mu}_t^X = -\hat{\mathbf{u}}_{t, \text{ML}}^X$ . Therefore,

$$\mathbb{E}_{p(\boldsymbol{\mu}_t^X)} \left[ \mathbf{u}_t^{X*} \mathbf{P} \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t) \right] = \mathbb{E}_{p(\boldsymbol{\mu}_t^X)} \left[ -\nabla_{\boldsymbol{\mu}_t^X} \cdot \mathbf{P} \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t) + \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t)^* \mathbf{P} \hat{\mathbf{u}}_{t, \text{ML}}^X \right] \quad (\text{S18})$$

where

$$\begin{aligned} \mathcal{L}_{\text{PCFM}}(\theta) &= \mathbb{E} \left[ \mathbf{u}_t^{X*} \mathbf{P} \mathbf{u}_t^X + \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t)^* \mathbf{P} \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t) + 2 \nabla_{\boldsymbol{\mu}_t^X} \cdot \mathbf{P} \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t) - 2 \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t)^* \mathbf{P} \hat{\mathbf{u}}_{t, \text{ML}}^X \right] \\ &= \mathbb{E} \left[ \left\| \mathbf{P}[\mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t) - \hat{\mathbf{u}}_{t, \text{ML}}^X] \right\|_2^2 + 2 \nabla_{\boldsymbol{\mu}_t^X} \cdot \mathbf{P} \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t) + \left\| \mathbf{P} \mathbf{u}_t^X \right\|_2^2 - \left\| \mathbf{P} \hat{\mathbf{u}}_{t, \text{ML}}^X \right\|_2^2 \right] \\ &= \mathbb{E}_{t, q(\mathbf{z}^X), p_t^X(\mathbf{x} | \mathbf{z}^X), p_t^Y(\mathbf{y} | \mathbf{z}^X)} \left[ \left\| \mathbf{P}[\mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t) - \hat{\mathbf{u}}_{t, \text{ML}}^X] \right\|_2^2 + 2 \nabla_{\boldsymbol{\mu}_t^X} \cdot \mathbf{P} \mathbf{v}_\theta^X(\boldsymbol{\mu}_t^X, t) \right] + \text{const.} \end{aligned} \quad (\text{S19})$$

Then, using Eq. (S9) and writing back  $\mathbf{v}_{\theta}^X(\boldsymbol{\mu}_t^X, t) = \mathbf{v}_{\theta}^X(\mathbf{A}^* \mathbf{y}, t)$ , we obtain by change of variables that  $\mathcal{L}_{\text{PCFM}}(\boldsymbol{\theta})$  can be transformed to the following expression up to a constant

$$\mathbb{E}_{t, q(\mathbf{z}^X), p_t^Y(\mathbf{y}|\mathbf{z}^X)} \left[ \left\| \mathbf{P}[\mathbf{v}_{\theta}^X(\mathbf{A}^* \mathbf{y}, t) - \hat{\mathbf{u}}_{t, \text{ML}}^X] \right\|_2^2 + \frac{2a_t}{a'_t} [(a'_t \sigma_0)^2 + (b'_t \sigma_1)^2] \nabla_{\mathbf{A}^* \mathbf{y}} \cdot \mathbf{P} \mathbf{v}_{\theta}^X(\mathbf{A}^* \mathbf{y}, t) \right]. \quad (\text{S20})$$

Finally, it concludes the proof by noting that

$$\mathbb{E}_{q(\mathbf{z}^X)} [p_t^Y(\mathbf{y} | \mathbf{z}^X)] = p_t^Y(\mathbf{y}) = \mathbb{E}_{q(\mathbf{z}^Y)} [p_t^Y(\mathbf{y} | \mathbf{z}^Y)], \quad (\text{S21})$$

where  $q(\mathbf{z}^Y) = q(\mathbf{y}_0) \mathbb{E}_{q(\mathbf{x}_1)} [p(\mathbf{y}_1 | \mathbf{x}_1)]$  is sampled by the MRI and Monte Carlo estimation.  $\square$

**Lemma 1.** *The  $\mathcal{Y}$ -space marginal vector field that generates the probability path  $p_t^Y$  takes the form*

$$\mathbf{u}_t^Y(\mathbf{y}) = \mathbf{A} \mathbf{v}_{\theta^*}^X(\mathbf{y}, t) - (a_t a'_t \sigma_0^2 + b_t b'_t \sigma_1^2) \nabla_{\mathbf{y}} \log p_t^Y(\mathbf{y}). \quad (\text{S22})$$

*Proof.* Eq. (16) shows that the  $\mathcal{Y}$ -space conditional vector field is

$$\mathbf{u}_t^Y(\mathbf{y} | \mathbf{z}^X) = \mathbf{A} \mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X) - (a_t a'_t \sigma_0^2 + b_t b'_t \sigma_1^2) \nabla_{\mathbf{y}} \log p_t^Y(\mathbf{y} | \mathbf{z}^X). \quad (\text{S23})$$

Therefore, by Proposition 1, the  $\mathcal{Y}$ -space marginal vector field takes the form

$$\begin{aligned} \mathbf{u}_t^Y(\mathbf{y}) &= \mathbb{E}_{q_t(\mathbf{z}^X | \mathbf{y})} [\mathbf{u}_t^Y(\mathbf{y} | \mathbf{z}^X)] \\ &= \mathbf{A} \mathbb{E}_{q_t(\mathbf{z}^X | \mathbf{y})} [\mathbf{u}_t^X(\mathbf{x} | \mathbf{z}^X)] - (a_t a'_t \sigma_0^2 + b_t b'_t \sigma_1^2) \mathbb{E}_{q_t(\mathbf{z}^X | \mathbf{y})} [\nabla_{\mathbf{y}} \log p_t^Y(\mathbf{y} | \mathbf{z}^X)] \\ &= \mathbf{A} \mathbf{v}_{\theta^*}^X(\mathbf{y}, t) - (a_t a'_t \sigma_0^2 + b_t b'_t \sigma_1^2) \nabla_{\mathbf{y}} \log p_t^Y(\mathbf{y}), \end{aligned} \quad (\text{S24})$$

where we have used the fact that

$$\begin{aligned} \mathbb{E}_{q_t(\mathbf{z}^X | \mathbf{y})} [\nabla_{\mathbf{y}} \log p_t^Y(\mathbf{y} | \mathbf{z}^X)] &= \int q_t(\mathbf{z}^X | \mathbf{y}) \frac{1}{p_t^Y(\mathbf{y} | \mathbf{z}^X)} \nabla_{\mathbf{y}} p_t^Y(\mathbf{y} | \mathbf{z}^X) d\mathbf{z}^X \\ &= \frac{1}{p_t^Y(\mathbf{y})} \int q(\mathbf{z}^X) \nabla_{\mathbf{y}} p_t^Y(\mathbf{y} | \mathbf{z}^X) d\mathbf{z}^X \\ &= \frac{1}{p_t^Y(\mathbf{y})} \nabla_{\mathbf{y}} \int q(\mathbf{z}^X) p_t^Y(\mathbf{y} | \mathbf{z}^X) d\mathbf{z}^X \\ &= \frac{1}{p_t^Y(\mathbf{y})} \nabla_{\mathbf{y}} p_t^Y(\mathbf{y}) \\ &= \nabla_{\mathbf{y}} \log p_t^Y(\mathbf{y}). \end{aligned} \quad (\text{S25})$$

$\square$

**Lemma 2.** *Note that  $p_1^Y(\mathbf{y}) = \int p_1(\mathbf{y} | \mathbf{x}) p_1^X(\mathbf{x}) d\mathbf{x} = \mathcal{CN}(\mathbf{y} | \mathbf{0}, 2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I}_{Cd})$ . Then,*

$$\mathbf{u}_t^Y(\mathbf{y}) = \frac{a'_t}{a_t} \mathbf{y} - b_t \left( b'_t - \frac{a'_t}{a_t} b_t \right) (2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I}_{Cd}) \nabla_{\mathbf{y}} \log p_t^Y(\mathbf{y}). \quad (\text{S26})$$

*Proof.* Taking  $\mathbf{y}_0$  as the conditioning variable, we have

$$\begin{aligned} \mathbf{u}_t^Y(\mathbf{y}) &= \mathbb{E}_{q_t(\mathbf{y}_0 | \mathbf{y})} [a'_t \mathbf{y}_0 + b'_t \mathbf{y}_1] \\ &= \mathbb{E}_{q_t(\mathbf{y}_0 | \mathbf{y})} \left[ a'_t \frac{\mathbf{y} - b_t \mathbf{y}_1}{a_t} + b'_t \mathbf{y}_1 \right] \\ &= \frac{a'_t}{a_t} \mathbf{y} + \left( b'_t - \frac{a'_t}{a_t} b_t \right) \mathbb{E}_{q_t(\mathbf{y}_0 | \mathbf{y})} [\mathbf{y}_1], \end{aligned} \quad (\text{S27})$$

where  $\mathbf{y}_1 = \frac{\mathbf{y} - a_t \mathbf{y}_0}{b_t}$ .

On the other hand, noting that  $p_t^Y(\mathbf{y} \mid \mathbf{y}_0) = \mathcal{CN}(\mathbf{y} \mid a_t \mathbf{y}_0, b_t^2(2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I}_{Cd}))$ , the score function can be written as

$$\begin{aligned}
 \nabla_{\mathbf{y}} \log p_t^Y(\mathbf{y}) &= \frac{1}{p_t^Y(\mathbf{y})} \nabla_{\mathbf{y}} p_t^Y(\mathbf{y}) \\
 &= \frac{1}{p_t^Y(\mathbf{y})} \nabla_{\mathbf{y}} \int p_t^Y(\mathbf{y} \mid \mathbf{y}_0) q(\mathbf{y}_0) d\mathbf{y}_0 \\
 &= \frac{1}{p_t^Y(\mathbf{y})} \int p_t^Y(\mathbf{y} \mid \mathbf{y}_0) q(\mathbf{y}_0) \nabla_{\mathbf{y}} \log p_t^Y(\mathbf{y} \mid \mathbf{y}_0) d\mathbf{y}_0 \\
 &= -\frac{1}{p_t^Y(\mathbf{y})} \int p_t^Y(\mathbf{y} \mid \mathbf{y}_0) q(\mathbf{y}_0) \frac{(2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I}_{Cd})^{-1}(\mathbf{y} - a_t \mathbf{y}_0)}{b_t^2} d\mathbf{y}_0 \\
 &= -\int q_t(\mathbf{y}_0 \mid \mathbf{y}) \frac{(2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I}_{Cd})^{-1} b_t \mathbf{y}_1}{b_t^2} d\mathbf{y}_0 \\
 &= -\frac{(2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I}_{Cd})^{-1}}{b_t} \mathbb{E}_{q_t(\mathbf{y}_0 \mid \mathbf{y})}[\mathbf{y}_1].
 \end{aligned} \tag{S28}$$

Combining Eq. (S27) and Eq. (S28) by canceling out  $\mathbb{E}_{q_t(\mathbf{y}_0 \mid \mathbf{y})}[\mathbf{y}_1]$  completes the proof.  $\square$

**Proposition 3 (Vector fields under projection).** For  $a_t = 1 - t$  and  $b_t = t$ , the  $\mathcal{Y}$ -space marginal vector field  $\mathbf{u}_t^Y(\mathbf{y})$  can be expressed by  $\mathbf{v}_{\theta^*}^X(\mathbf{y}, t)$  as

$$\mathbf{u}_t^Y(\mathbf{y}) = \mathbf{A} \mathbf{v}_{\theta^*}^X(\mathbf{y}, t) - \frac{c_t}{1-t} [(c_t + \sigma_1^2) \mathbf{I}_{Cd} + 2\mathbf{A}\mathbf{A}^*]^{-1} [(1-t) \mathbf{A} \mathbf{v}_{\theta^*}^X(\mathbf{y}, t) + \mathbf{y}], \tag{S29}$$

where  $c_t \triangleq (1-t) \left( \frac{1-t}{t} \sigma_0^2 - \sigma_1^2 \right)$ . In addition, left-multiplying both sides with  $\mathbf{A}^*$  gives the more computationally friendly formula when  $Cd > D$ :

$$\mathbf{A}^* \mathbf{u}_t^Y(\mathbf{y}) = \mathbf{A}^* \mathbf{A} \mathbf{v}_{\theta^*}^X(\mathbf{y}, t) - \frac{c_t}{1-t} [(c_t + \sigma_1^2) \mathbf{I}_D + 2\mathbf{A}^* \mathbf{A}]^{-1} \mathbf{A}^* [(1-t) \mathbf{A} \mathbf{v}_{\theta^*}^X(\mathbf{y}, t) + \mathbf{y}]. \tag{S30}$$

*Proof.* For  $a_t = 1 - t$  and  $b_t = t$ , Lemma 2 indicates

$$\nabla_{\mathbf{y}} \log p_t^Y(\mathbf{y}) = -\frac{1}{t} (2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I}_{Cd})^{-1} [\mathbf{y} + (1-t) \mathbf{u}_t^Y(\mathbf{y})]. \tag{S31}$$

Substituting this into Lemma 1 gives the equation

$$\mathbf{u}_t^Y(\mathbf{y}) = \mathbf{A} \mathbf{v}_{\theta^*}^X(\mathbf{y}, t) - \left( \frac{1-t}{t} \sigma_0^2 - \sigma_1^2 \right) (2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I}_{Cd})^{-1} [\mathbf{y} + (1-t) \mathbf{u}_t^Y(\mathbf{y})]. \tag{S32}$$

Denoting  $\mathbf{u} \triangleq \mathbf{u}_t^Y(\mathbf{y})$  and  $\mathbf{v} \triangleq \mathbf{v}_{\theta^*}^X(\mathbf{y}, t)$ , then solving for  $\mathbf{u}$  in the above equation gives

$$\mathbf{u} = (\mathbf{I} + c(2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I}_{Cd})^{-1})^{-1} \left( \mathbf{A} \mathbf{v} - \frac{c}{1-t} (2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I}_{Cd})^{-1} \mathbf{y} \right), \tag{S33}$$

where  $c \triangleq \left( \frac{1-t}{t} \sigma_0^2 - \sigma_1^2 \right) (1-t)$ .

By the Woodbury matrix identity  $(\mathbf{A} + \mathbf{UCV})^{-1} = \mathbf{A}^{-1} - \mathbf{A}^{-1} \mathbf{U} (\mathbf{C}^{-1} + \mathbf{V} \mathbf{A}^{-1} \mathbf{U}^{-1})^{-1} \mathbf{V} \mathbf{A}^{-1}$ , we have

$$(\mathbf{I} + c(2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I})^{-1})^{-1} = \mathbf{I} - c(c\mathbf{I} + 2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I})^{-1}, \tag{S34}$$

and note that

$$(\mathbf{I} + c(2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I})^{-1})^{-1} (2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I})^{-1} = (c\mathbf{I} + 2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I})^{-1}. \tag{S35}$$

Therefore,

$$\begin{aligned}
 \mathbf{u} &= \left[ \mathbf{I} - c(c\mathbf{I} + 2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I})^{-1} \right] \mathbf{A} \mathbf{v} - \frac{c}{1-t} (c\mathbf{I} + 2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I})^{-1} \mathbf{y} \\
 &= \mathbf{A} \mathbf{v} - c(c\mathbf{I} + 2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I})^{-1} \left[ \mathbf{A} \mathbf{v} + \frac{1}{1-t} \mathbf{y} \right] \\
 &= \mathbf{A} \mathbf{v} - \frac{c}{1-t} (c\mathbf{I} + 2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I})^{-1} [(1-t) \mathbf{A} \mathbf{v} + \mathbf{y}],
 \end{aligned} \tag{S36}$$

and by the identity  $\mathbf{A}^*(c\mathbf{I}_{Cd} + 2\mathbf{A}\mathbf{A}^* + \sigma_1^2 \mathbf{I}_{Cd})^{-1} = (c\mathbf{I}_D + 2\mathbf{A}^* \mathbf{A} + \sigma_1^2 \mathbf{I}_D)^{-1} \mathbf{A}^*$ ,

$$\mathbf{A}^* \mathbf{u} = \mathbf{A}^* \mathbf{A} \mathbf{v} - \frac{c}{1-t} (c\mathbf{I}_D + 2\mathbf{A}^* \mathbf{A} + \sigma_1^2 \mathbf{I})^{-1} \mathbf{A}^* [(1-t) \mathbf{A} \mathbf{v} + \mathbf{y}]. \tag{S37}$$

$\square$

**Algorithm 2:** PnP Cyclic Integration with PCFM for Noisy Measurements

**Input:** k-space measurement  $\mathbf{y}_0$ , pretrained optimal solution to PCFM  $\mathbf{v}_{\theta^*}^X(\cdot, t)$ , number of time steps  $T$ , adaptive step size  $\gamma_t$ .

**Output:** Reconstructed image  $\mathbf{x}_0$  of  $\mathbf{y}_0$ .

```

1 for  $t = 0, \dots, (T-1)/T$  do
2    $\mathbf{y}_{t+1/T} \leftarrow \mathbf{y}_t + \frac{1}{T} \mathbf{u}_t^Y(\mathbf{y})$  ▷ Forward integration using Proposition 3
3   Sample  $\mathbf{x}_1 \sim p_1(\mathbf{x} \mid \mathbf{y})$ . ▷ Posterior sampling with Eq. (22)
4   for  $t \in \{1, \dots, 1/T\}$  do
5      $\tilde{\mathbf{y}}_t \leftarrow a_t \mathbf{y}_0 + b_t \mathbf{y}_1$ 
6      $\tilde{\mathbf{x}}_0 \leftarrow \mathbf{x}_t - t \mathbf{v}_{\theta^*}^X(\tilde{\mathbf{y}}_t, t)$  ▷ PnP-Flow denoising step
7      $\mathbf{x}_0 \leftarrow \tilde{\mathbf{x}}_0 - \gamma_t \mathbf{A}^*(\mathbf{A} \tilde{\mathbf{x}}_0 - \tilde{\mathbf{y}}_0)$  ▷ Gradient step
8      $\mathbf{x}_{t-1/T} \leftarrow a_{t-1/T} \mathbf{x}_0 + b_{t-1/T} \mathbf{x}_1$  ▷ Interpolation step
9 return  $\mathbf{x}_0$ .
```

## B PLUG-AND-PLAY CYCLIC INTEGRATION WITH PCFM FOR NOISY MEASUREMENTS

The plug-and-play (PnP) framework (Venkatakrishnan et al., 2013; Fang et al., 2024) utilizes off-the-shelf denoising techniques to address the general inverse problem while optimizing the objective function

$$\min_{\mathbf{x}} \left\{ \frac{1}{2} \|\mathbf{y} - \mathcal{A}(\mathbf{x})\|_2^2 + R(\mathbf{x}) \right\}, \quad (\text{S38})$$

where  $R(\mathbf{x})$  serves as a regularizer, encouraging solutions that are probable under the prior distribution of  $\mathbf{x}$ . PnP substitutes the explicit solution of the proximal operator

$$\text{prox}_R(\mathbf{y}) \triangleq \arg \min_{\mathbf{x}} \left\{ \frac{1}{2} \|\mathbf{y} - \mathbf{x}\|_2^2 + R(\mathbf{x}) \right\}, \quad (\text{S39})$$

which is utilized to optimize the objective through methods such as proximal gradient descent (PGD) (Beck, 2017), half quadratic splitting (HQS) (Geman & Yang, 1995), and alternating direction methods of multipliers (ADMM) (Boyd et al., 2011). In situations involving noisy k-space data, we utilize the PGD and PnP-Flow frameworks (Martin et al., 2025). These frameworks alternate between a denoising step using a pretrained flow matching model, aiming to approximate the proximal operator’s solution, and a gradient step to promote data consistency. Alg. 2 presents the discrete-time algorithm implementing the PnP-based cyclic integration with the proposed PCFM framework.

## C NUMERICAL METHODS

### C.1 CONJUGATE GRADIENT

Conjugate gradient (CG) is a method for solving a linear system of equations  $\mathbf{A}\mathbf{x} = \mathbf{b}$ , when the matrix  $\mathbf{A}$  is symmetric (or Hermitian) positive-definite (SPD) and very large (often sparse). Direct methods like Gaussian elimination are impractical due to computational cost and memory requirements. CG reframes the problem of solving a linear system as an optimization problem. The solution to  $\mathbf{A}\mathbf{x} = \mathbf{b}$  is precisely the vector  $\mathbf{x}$  that minimizes the quadratic form:

$$\phi(\mathbf{x}) = \frac{1}{2} \mathbf{x}^* \mathbf{A} \mathbf{x} - \mathbf{x}^* \mathbf{b}. \quad (\text{S40})$$

An intuitive way to find this minimum is to take the steepest descent, where one repeatedly steps in the direction of the negative gradient  $-\nabla \phi(\mathbf{x}) = \mathbf{b} - \mathbf{A}\mathbf{x}$ . However, the steepest descent often performs poorly, taking many small zigzagging steps to reach the minimum.

CG dramatically improves upon this by choosing a sequence of search directions that are “smarter”. Instead of using the residual at each step, it generates a set of search directions  $\{\mathbf{p}_0, \mathbf{p}_1, \dots, \mathbf{p}_{k-1}\}$  that are mutually  $\mathbf{A}$ -orthogonal, i.e.,  $\mathbf{p}_i^* \mathbf{A} \mathbf{p}_j = 0$  for  $i \neq j$ . This property is crucial: minimizing the

**Algorithm 3:** Conjugate Gradient (CG)

---

**Input:** A symmetric (or Hermitian) matrix  $\mathbf{A}$ , a vector  $\mathbf{b}$ , an initial guess  $\mathbf{x}_0$ , a maximum number of iterations  $k$ , and a tolerance  $\epsilon$ .

**Output:** An approximate solution  $\mathbf{x}$  to  $\mathbf{Ax} = \mathbf{b}$ .

```

1 Set  $\mathbf{r}_0 \leftarrow \mathbf{b} - \mathbf{Ax}_0$  and  $\mathbf{p}_0 \leftarrow \mathbf{r}_0$ . ▷ Initialization
2 for  $j = 1, \dots, k$  do
3    $\mathbf{v}_j \leftarrow \mathbf{Ap}_j$ 
4    $\alpha_j \leftarrow \frac{\mathbf{r}_j^* \mathbf{r}_j}{\mathbf{p}_j^* \mathbf{v}_j}$  ▷ Compute the step size
5    $\mathbf{x}_{j+1} \leftarrow \mathbf{x}_j + \alpha_j \mathbf{p}_j$  ▷ Update solution
6    $\mathbf{r}_{j+1} \leftarrow \mathbf{r}_j - \alpha_j \mathbf{v}_j$  ▷ Update residual
7   if  $\|\mathbf{r}_{j+1}\|_2 < \epsilon$  then
8     break
9    $\beta_j \leftarrow \frac{\mathbf{r}_{j+1}^* \mathbf{r}_{j+1}}{\mathbf{r}_j^* \mathbf{r}_j}$  ▷ Calculate the improvement factor
10   $\mathbf{p}_{j+1} \leftarrow \mathbf{r}_{j+1} + \beta_j \mathbf{p}_j$  ▷ Update search direction
11 return  $\mathbf{x}_{j+1}$ .
```

---

quadratic function along a new search direction  $\mathbf{p}_k$  does not compromise the minimization that has already been achieved in the previous directions. At each iteration  $k$ , CG finds the optimal solution  $\mathbf{x}_k$  within the affine Krylov subspace  $\mathbf{x}_0 + \mathcal{K}_k(\mathbf{A}, \mathbf{r}_0)$ , where  $\mathbf{r}_0 \triangleq \mathbf{b} - \mathbf{Ax}_0$  is the initial residual. This means that after  $k$  steps, CG has found the best possible solution that can be formed by a linear combination of  $\{\mathbf{r}_0, \mathbf{Ar}_0, \dots, \mathbf{A}^{k-1}\mathbf{r}_0\}$ . This guarantees convergence for an  $N \times N$  matrix in at most  $N$  steps, though in practice, a good approximation is often found in far fewer iterations. Alg. 3 provides the algorithm in detail.

## D DETAILS OF THE COMPARED BASELINES

We benchmark our method against three types of baseline approaches.

(A) Supervised methods that require fully sampled MRIs during training:

- **MoDL** (Aggarwal et al., 2018) is a model-based end-to-end network that unrolls traditional optimization procedure by regarding the CNN as a regularizer. We use 10 unrolling iterations of the network. Training is performed by minimizing the L2 loss between the network output and the ground truth.
- **DDNM<sup>+</sup>** (Wang et al., 2023) is a diffusion model-based method for inverse problems solving. Training is performed by denoising on fully sampled images with the DDPM framework (Ho et al., 2020), whereas the inference is implemented by alternating between the DDIM sampling steps (Song et al., 2021) and the range-null decomposition for enforcing measurement consistency. Training is performed by the noise prediction objective with the cosine noise schedule of iDDPM (Nichol & Dhariwal, 2021). We use the same network architecture and training strategy as our proposed approach for this method.
- **OT-ODE** (Poke et al., 2024) estimates the posterior vector field by combining the original vector field learned from fully sampled images with the likelihood score approximated by the IIGDM estimation (Song et al., 2023). Training is performed by optimizing the original CFM objective (Lipman et al., 2023). We use the same network architecture and training strategy as our proposed approach for this method.
- **PnP-Flow** (Martin et al., 2025) integrates the Plug-and-Play framework with flow matching, which alternates between gradient descent steps for measurement consistency, reprojections onto the learned flow path, and denoising by the pre-trained vector field. Training is performed by optimizing the original CFM objective (Lipman et al., 2023). We use the same network architecture and training strategy as our proposed approach for this method. We set the hyperparameter  $\gamma_t = t^\alpha$  with  $\alpha = 0.1$ .

(B) Self-supervised methods that learn to reconstruct by additional subsampling of the available k-space measurements.

- **SSDU** (Yaman et al., 2020) proposes to split the available k-space measurements into two disjoint subsets and then train a model-based reconstruction network to recover one of the subsets from the other. We use the VarNet (Hammernik et al., 2018) with 10 unrolling iterations as the network backbone. Training is achieved by minimizing the L2 loss.
- **Weighted SSDU** (Millard & Chiew, 2023) improves upon the SSDU framework by using a subsampling mask of the same distribution as the original mask and re-weighting the L2 loss. We use the VarNet (Hammernik et al., 2018) with 10 unrolling iterations as the network backbone.
- **Robust SSDU** (Millard & Chiew, 2024) provably recovers clean images from noisy, under-sampled training data by simultaneously estimating missing k-space samples and denoising the available samples. We use the VarNet (Hammernik et al., 2018) with 10 unrolling iterations as the network backbone.

(C) Unsupervised methods that learn to reconstruct using all the available k-space measurements.

- **REI** (Chen et al., 2022) achieves unsupervised reconstruction by combining the k-space SURE-based loss (Stein, 1981) for measurement consistency and the equivariant imaging framework which builds on the group invariance assumption of the signal space (Chen et al., 2021; Tachella et al., 2023). We use the VarNet (Hammernik et al., 2018) with 10 unrolling iterations as the network backbone.
- **MOI** (Tachella et al., 2022) leverages the randomness in the imaging operator and proposes an unsupervised loss that ensures consistency across all operators. We use the MoDL (Aggarwal et al., 2018) architecture with 10 unrolling iterations as the network backbone.
- **ENSURE** (Aggarwal et al., 2022) also leverages the randomness in the forward operators and provides an unbiased estimate of the true mean squared error without fully sampled images. We use the MoDL (Aggarwal et al., 2018) architecture with 10 unrolling iterations as the network backbone. However, it uses an inaccurate numerical strategy to approximate the loss function in the multi-coil scenario.
- **GTF<sup>2</sup>M** (Luo et al., 2025) proposes a ground-truth-free flow matching framework for single-coil MRI. Reconstruction for multi-coil MRI is achieved by coil-wise reconstruction followed by SENSE-based combination. We use the same network architecture and training strategy as our proposed approach for this method. Nevertheless, the performance of this method is suboptimal as the prior is learned from single-coil k-space measurements instead of the combined forward operator of parallel MRI.

## E ADDITIONAL RESULTS

### E.1 INFERENCE TIME

Table S1 illustrates the average inference time of the various methods evaluated for the reconstruction of a single image from the CMRxRecon 2023 dataset. It is evident that among the generative model-based techniques, our method demonstrates the fastest inference time.

### E.2 ABLATION STUDY ON BACKWARD SAMPLING

Table S2 displays the results of the ablation study examining various backward sampling strategies applied to the original initial noiseless CMRxRecon 2023 dataset. It can be noted that the PnP-based backward sampling technique is less effective compared to Alg. 1 when tested on noiseless data.

### E.3 ABLATION STUDY ON FORWARD SAMPLING

If the forward integration is omitted, the unconditional distribution  $p(\mathbf{y}_1 | \mathbf{y}_0)$  will be assumed to be  $\mathcal{CN}(\mathbf{0}, \mathbf{I}_d)$ . Table S3 displays results from the ablation study investigating the impact of incorporating the proposed forward integration steps. It is evident that incorporating forward integration steps markedly enhances reconstruction performance.

Table S1: Average inference time of the compared methods for reconstructing one image from the CMRxRecon 2023 dataset. The inference time is calculated on an NVIDIA A5000 GPU using a batch size of 4.

| Training Supervision | Reconstruction Method | Time (ms) | NFEs ↓ |
|----------------------|-----------------------|-----------|--------|
| Supervised           | MoDL                  | 58        | 1      |
|                      | DDNM <sup>+</sup>     | 1000      | 100    |
|                      | OT-ODE                | 1750      | 100    |
|                      | PnP-Flow              | 938       | 100    |
| Self-supervised      | SSDU                  | 19        | 1      |
|                      | Weighted SSDU         | 19        | 1      |
|                      | Robust SSDU           | 19        | 1      |
| Unsupervised         | REI                   | 19        | 1      |
|                      | MOI                   | 58        | 1      |
|                      | ENSURE                | 19        | 1      |
|                      | GTF <sup>2</sup> M    | 1563      | 20     |
|                      | PCFM (Ours)           | 500       | 20     |

Table S2: Ablation study on the backward integration strategy on the CMRxRecon 2023 cardiac T1/T2 mapping MRI (noiseless). The difference in metrics is statistically significant between the two strategies by the two-sided paired  $t$ -test ( $p < 0.05$ ).

| Reconstruction Method | SSIM ↑            |                   | PSNR ↑           |                  | NFEs ↓ |
|-----------------------|-------------------|-------------------|------------------|------------------|--------|
|                       | 4×                | 8×                | 4×               | 8×               |        |
| Zero-Filled           | $0.769 \pm 0.057$ | $0.747 \pm 0.056$ | $27.01 \pm 1.97$ | $26.11 \pm 1.90$ | N/A    |
| PnP-PCFM              | $0.808 \pm 0.057$ | $0.893 \pm 0.038$ | $28.82 \pm 2.19$ | $33.16 \pm 2.82$ | 20     |
| PCFM                  | $0.994 \pm 0.008$ | $0.974 \pm 0.016$ | $52.39 \pm 9.81$ | $40.50 \pm 3.84$ | 20     |

#### E.4 NOISY DATA

Table S4 displays the quantitative outcomes on the noise-free test data from the CMRxRecon 2023 dataset. A notable observation is that the model retains its performance as though the training data were free of noise, suggesting that PCFM effectively learns the prior distribution of fully sampled MRI even when trained on noisy and undersampled inputs.

Fig. S2 illustrates the reconstruction outcomes from cardiac MRI when subjected to additive Gaussian noise at different levels. It can be noted that with an increase in noise level, the reconstruction is more susceptible to contamination by measurement noise.

Table S3: Ablation study on the forward sampling strategy on the CMRxRecon 2023 cardiac T1/T2 mapping MRI (noiseless). The difference in metrics is statistically significant between the two strategies by the two-sided paired  $t$ -test ( $p < 0.05$ ).

| Reconstruction Method | SSIM ↑            |                   | PSNR ↑           |                  | NFEs ↓ |
|-----------------------|-------------------|-------------------|------------------|------------------|--------|
|                       | 4×                | 8×                | 4×               | 8×               |        |
| Zero-Filled           | $0.769 \pm 0.057$ | $0.747 \pm 0.056$ | $27.01 \pm 1.97$ | $26.11 \pm 1.90$ | N/A    |
| PCFM w/o forward      | $0.988 \pm 0.012$ | $0.960 \pm 0.021$ | $46.21 \pm 6.00$ | $38.12 \pm 3.51$ | 10     |
| PCFM                  | $0.994 \pm 0.008$ | $0.974 \pm 0.016$ | $52.39 \pm 9.81$ | $40.50 \pm 3.84$ | 20     |

Table S4: Qualitative results of  $4\times$  and  $8\times$  parallel MRI reconstruction of CMRxRecon 2023 cardiac T1/T2 quantitative mapping images using PCFM trained on noisy data while tested on clean data.

| Training Noise level | Reconstruction Method | SSIM $\uparrow$   |                   | PSNR $\uparrow$  |                  | NFEs $\downarrow$ |
|----------------------|-----------------------|-------------------|-------------------|------------------|------------------|-------------------|
|                      |                       | $4\times$         | $8\times$         | $4\times$        | $8\times$        |                   |
| $\sigma_0 = 0.05$    | PCFM (Ours)           | $0.995 \pm 0.005$ | $0.974 \pm 0.016$ | $51.34 \pm 6.21$ | $40.51 \pm 3.84$ | 20                |
| $\sigma_0 = 0.1$     | PCFM (Ours)           | $0.995 \pm 0.005$ | $0.974 \pm 0.016$ | $51.34 \pm 6.19$ | $40.50 \pm 3.83$ | 20                |
| $\sigma_0 = 0.2$     | PCFM (Ours)           | $0.995 \pm 0.005$ | $0.974 \pm 0.016$ | $51.30 \pm 6.21$ | $40.55 \pm 3.83$ | 20                |

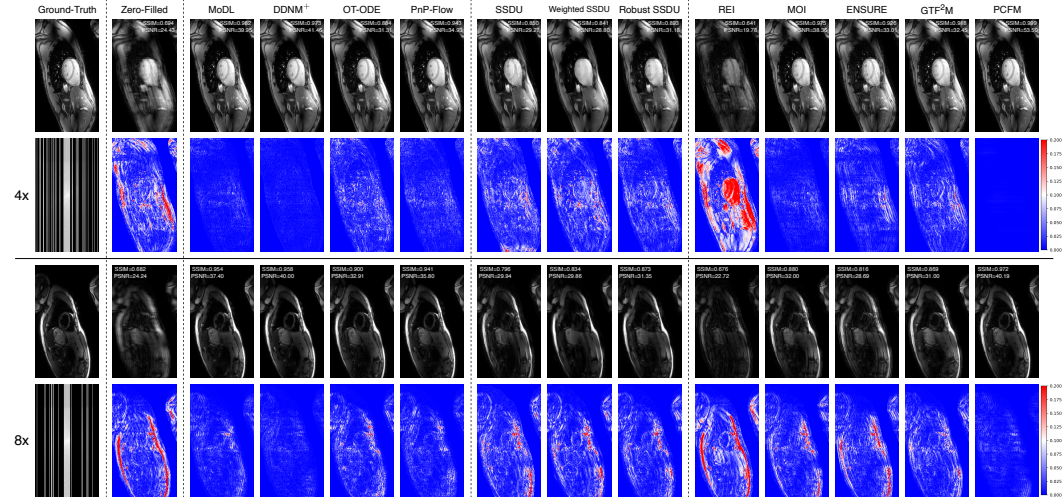


Figure S1: Visualization of reconstruction on two test samples of  $4\times$  and  $8\times$  accelerated multi-coil cardiac MRI from the compared methods. The k-space are presented in log-scale absolute values. The error maps are presented in values relative to the peak intensity in the ground-truth image.

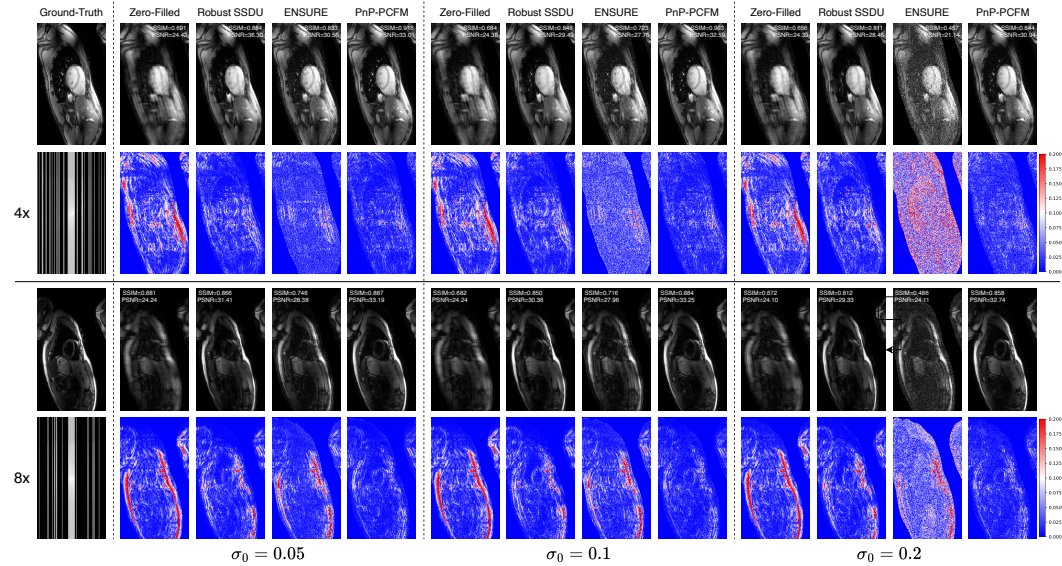


Figure S2: Visualization of reconstruction on two test samples of  $4\times$  and  $8\times$  accelerated multi-coil cardiac MRI with additive Gaussian noise of various scales.