

Stabilizing Physics-Informed Consistency Models via Structure-Preserving Training

Che-Chia Chang
Institute of Artificial Intelligence
Innovation
National Yang Ming Chiao Tung
University
Hsinchu, Taiwan
ccchang.c@nycu.edu.tw

Chen-Yang Dai
Department of Applied Mathematics
National Yang Ming Chiao Tung
University
Hsinchu, Taiwan
cydai.sc12@nycu.edu.tw

Te-Sheng Lin
Department of Applied Mathematics
National Yang Ming Chiao Tung
University
Hsinchu, Taiwan
National Center for Theoretical
Sciences
National Taiwan University
Taipei, Taiwan
teshenglin@nycu.edu.tw

Ming-Chih Lai
Department of Applied Mathematics
National Yang Ming Chiao Tung
University
Hsinchu, Taiwan
mclai@math.nctu.edu.tw

Chieh-Hsin Lai
Department of Applied Mathematics
National Yang Ming Chiao Tung
University
Hsinchu, Taiwan
chiehhsinlai@gmail.com

Abstract

We propose a physics-informed consistency modeling framework for solving partial differential equations (PDEs) via fast, few-step generative inference. We identify a key stability challenge in physics-constrained consistency training, where PDE residuals can drive the model toward trivial or degenerate solutions, degrading the learned data distribution. To address this, we introduce a structure-preserving two-stage training strategy that decouples distribution learning from physics enforcement by freezing the coefficient decoder during physics-informed fine-tuning. We further propose a two-step residual objective that enforces physical consistency on refined, structurally valid generative trajectories rather than noisy single-step predictions. The resulting framework enables stable, high-fidelity inference for both unconditional generation and forward problems. We demonstrate that forward solutions can be obtained via a projection-based zero-shot inpainting procedure, achieving consistent accuracy of diffusion baselines with orders of magnitude reduction in computational cost.

CCS Concepts

• **Computing methodologies** → **Neural networks**; *Learning latent representations*; *Modeling methodologies*; • **Applied computing** → *Physics*.

Keywords

physics-informed, scientific machine learning, partial differential equations, PDE Solving, consistency model, generative modeling



This work is licensed under a Creative Commons Attribution 4.0 International License.
KDD '26, Jeju Island, Republic of Korea
© 2026 Copyright held by the owner/author(s).
ACM ISBN 979-8-4007-2259-2/2026/08
<https://doi.org/10.1145/3770855.3819059>

ACM Reference Format:

Che-Chia Chang, Chen-Yang Dai, Te-Sheng Lin, Ming-Chih Lai, and Chieh-Hsin Lai. 2026. Stabilizing Physics-Informed Consistency Models via Structure-Preserving Training. In *Proceedings of the 32nd ACM SIGKDD Conference on Knowledge Discovery and Data Mining V.2 (KDD '26)*, August 09–13, 2026, Jeju Island, Republic of Korea. ACM, New York, NY, USA, 12 pages. <https://doi.org/10.1145/3770855.3819059>

1 Introduction

Diffusion Models [5, 10, 20, 22, 23] have achieved remarkable success in high-fidelity generation of images and audio. Motivated by their ability to model complex, high-dimensional distributions, recent works have adapted diffusion models to the domain of scientific machine learning [7, 19]. Unlike traditional numerical solvers or Physics-Informed Neural Networks (PINNs) [18], which determine a single solution to a given PDE, these generative models learn the full distribution of possible solutions. This capability is crucial for ill-posed inverse problems, such as reconstructing complex physical fields from sparse observations [6], where a single deterministic prediction fails to capture the multiplicity of valid states.

However, despite their expressive power, diffusion models suffer from iterative sampling procedures that require hundreds or even thousands of neural network evaluations to generate a single sample. This limits their practicality for scientific applications that demand real-time or large-scale simulations. To overcome this bottleneck, Consistency Models (CMs) [21] have been proposed as a fast alternative, enabling one- or few-step generation from noise to data.

While Consistency Models offer a promising path toward real-time scientific simulation, adapting them to strictly enforce physical laws introduces additional optimization challenges. Unlike standard image generation, where perceptual quality is the primary

metric, scientific solutions must adhere to rigorous governing equations. Empirically, we observe that directly incorporating Partial Differential Equation (PDE) residuals into the consistency objective often leads to unstable training dynamics and suboptimal convergence. One common difficulty in this setting is that physics losses may admit low-complexity or degenerate solutions, such as near-zero or overly smooth fields, that yield small residual values but fail to represent meaningful physical states. In practice, this often manifests as severe mode collapse, where consistency training converges to a narrow subset of admissible solutions despite satisfying the imposed physical constraints. This mismatch between physical constraints and the self-consistency objective motivates the need for specialized training and inference strategies to obtain both distributional fidelity and physical correctness.

In this work, we propose sCM-PINN, a physics-informed consistency modeling framework for efficient inference in forward and conditional PDE problems. Our approach introduces a two-stage training strategy that decouples distribution learning from physics-constrained fine-tuning, thereby stabilizing optimization in the presence of sensitive residual losses. To prevent the physical constraints from degrading the learned data manifold, we partition the model output into coefficient and solution channels, freezing the coefficient decoder during the physics-informed phase. Furthermore, to bridge the gap between consistency training and physical validity, we introduce a two-step residual-based constraint that enforces PDE consistency across refined trajectories rather than noisy, single-step predictions. Together, these design choices enable fast, reliable inference for both unconditional sampling and forward problem solving.

Our main contributions are:

- **Two-Stage Training Framework:** A stabilized optimization pipeline for physics-informed consistency models that prevents gradient interference between data and physics objectives.
- **Channel-Partitioned Architecture:** A strategy to preserve the integrity of the learned data distribution by isolating and freezing the coefficient manifold.
- **Two-Step PDE Residual Constraint:** A robust physical supervision mechanism that aligns the consistency sampling trajectory with the underlying PDE operator.
- **Projection-Based Inference:** A zero-shot inpainting formulation for forward PDE problems that bypasses the need for expensive test-time gradient descent.
- **Physical Evaluation Metric:** The introduction of the relative H^1 error metric to rigorously quantify the alignment of generated samples with the PDE, ensuring consistency in both solution values and their spatial derivatives.

Finally, for reproducibility and future research, we release our complete codebase at: <https://github.com/twMisc/sCM-PINN>.

2 Preliminaries and Related Work

2.1 Generative Modeling of PDE Solutions

Conventional learning-based approaches for scientific computing, such as Neural Operators [12, 16], typically formulate PDE solving as learning a deterministic mapping from coefficient fields to solution fields. This formulation is highly effective for well-posed

forward problems, but becomes fundamentally limited in inverse or ill-posed settings, where solutions may be non-unique, and observations are sparse or incomplete.

To address these limitations, recent work has reformulated PDE solving as a *generative modeling* problem. Methods such as DiffusionPDE [6] aim to learn the joint distribution $p(\mathbf{a}, \mathbf{u})$ over physical coefficients $\mathbf{a}$ and state variables $\mathbf{u}$. Let $\mathbf{x} = [\mathbf{a}, \mathbf{u}] \in \mathbb{R}^D$ denote the concatenated system state. By learning the score function $\nabla_{\mathbf{x}} \log p(\mathbf{x})$, these models unify forward and inverse problems within a single conditional sampling framework: forward prediction corresponds to sampling from $p(\mathbf{u} | \mathbf{a})$, while inverse inference corresponds to sampling from $p(\mathbf{a} | \mathbf{u})$. However, such score-based approaches typically rely on iterative sampling procedures, leading to high inference costs in high-dimensional PDE settings.

2.2 Consistency Models for Fast Sampling

Consistency Models [21] are a class of generative models designed to overcome the inference inefficiency of diffusion-based sampling. Instead of iteratively integrating a reverse-time Stochastic differential equation or Probability Flow ODE (PF-ODE) [23], CMs learn a continuous consistency function $f_{\theta}(\mathbf{x}_t, t)$ that maps any point $\mathbf{x}_t$ along the PF-ODE trajectory directly to its corresponding data sample $\mathbf{x}_0$. The training objective enforces a *self-consistency* condition,

$$f_{\theta}(\mathbf{x}_t, t) = f_{\theta}(\mathbf{x}_{t'}, t'), \quad \forall t, t' \in [0, T], \quad (1)$$

ensuring that predictions remain invariant across the trajectory. By distilling the iterative integration process into a single mapping, CMs enable one-step or few-step generative inference. In this work, we adopt the continuous-time formulation (sCM) [15], which eliminates discretization error present in discrete-time consistency training and improves training stability.

2.3 Physics-Informed Generative Models

While consistency models (CMs) provide an efficient framework for few-step generative inference, they lack the inherent inductive biases required to satisfy the underlying governing equations of physical systems. Bridging this gap necessitates the integration of PDE constraints directly into the CM training and sampling trajectories. Within the broader literature of generative modeling for PDEs, this integration typically follows two distinct paradigms: inference-time guidance, which steers the sampling process toward physical validity, and training-time supervision, which embeds physical laws into the model's learned parameters via residual-based loss functions.

Inference-Time Guidance. Inference-time guidance methods, such as Diffusion Posterior Sampling (DPS) [4], enforce physical constraints during the sampling process by modifying the score function with gradients derived from the PDE residual operator $\mathcal{R}(\cdot)$. A typical formulation is

$$\nabla_{\mathbf{x}_t} \log p(\mathbf{x}_t | \mathcal{R} = 0) \approx \nabla_{\mathbf{x}_t} \log p(\mathbf{x}_t) - \eta \nabla_{\mathbf{x}_t} \|\mathcal{R}(\hat{\mathbf{x}}_0(\mathbf{x}_t))\|^2. \quad (2)$$

While flexible and model-agnostic, this approach incurs substantial computational overhead, as it requires backpropagation through

both the neural network and the differential operator at each sampling step. Moreover, gradients of PDE residuals can be poorly conditioned, often necessitating careful tuning of guidance strengths across different problem settings.

Training-Time Supervision. An alternative paradigm incorporates physics constraints directly into the training objective. Physics-Informed Diffusion Models (PIDMs) [2] optimize a composite loss that combines the denoising objective with a PDE residual penalty:

$$\mathcal{L}_{\text{PIDM}}(\theta) = \mathbb{E}_{t, \mathbf{x}_0} \left[\ell_{\text{denoise}}(\theta) + \lambda \left\| \mathcal{R}(\hat{\mathbf{x}}_0^\theta) \right\|^2 \right]. \quad (3)$$

By amortizing physics enforcement into training, this strategy enables fast, physics-consistent generation at inference time. However, extending this paradigm to consistency models presents additional challenges. As discussed in Section 3, the consistency objective can be sensitive to gradient variance introduced by PDE residual terms, motivating the stabilized training strategy proposed in this work.

3 Methodology

We present a physics-informed consistency training framework designed to bridge the gap between fast generative inference and rigorous physical validity for forward PDE problems. Our approach builds upon the sCM framework [15]. In incorporating physics-based constraints, we identify a common optimization failure mode in which the model collapses toward degenerate solutions that minimize the PDE residual but fail to preserve the target data distribution. To address this issue, we propose a training strategy that combines a two-stage optimization protocol with a structure-preserving channel partition.

3.1 Continuous-Time Consistency Parameterization

We adopt the TrigFlow parameterization introduced in the sCM framework [15]. This formulation learns a consistency function $f_\theta(\mathbf{x}_t, t)$ that maps points along the probability flow trajectory directly to the corresponding clean data sample $\mathbf{x}_0$.

Let the state $\mathbf{x} = [\mathbf{a}, \mathbf{u}]$ represent the concatenation of the coefficient field $\mathbf{a}$ and the solution field $\mathbf{u}$. For a state $\mathbf{x}_t$ at time $t \in [0, \pi/2]$, the model is defined as:

$$f_\theta(\mathbf{x}_t, t) = \cos(t)\mathbf{x}_t - \sin(t)\sigma_d F_\theta\left(\frac{\mathbf{x}_t}{\sigma_d}, t\right), \quad (4)$$

where F_θ is a neural network parameterizing the normalized velocity field and σ_d denotes the data standard deviation. This parameterization smoothly interpolates between the data distribution (at $t = 0$) and a Gaussian prior (at $t = \pi/2$), and serves as a numerically stable backbone for incorporating physical constraints.

3.2 Stabilizing Physics-Constrained Consistency Training

Integrating PDE constraints into consistency models introduces significant optimization challenges. In particular, directly optimizing a composite objective from random initialization often leads to mode collapse or convergence to trivial solutions. As illustrated in Figure 1 (Left), a model trained directly on the physics objective

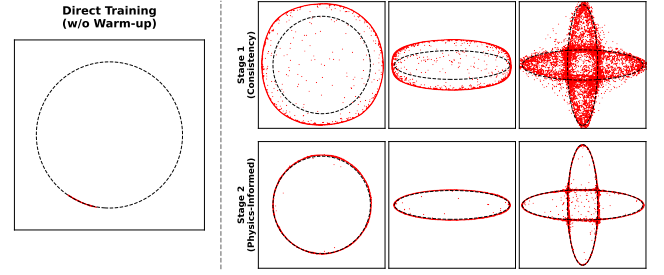


Figure 1: Necessity of Two-Stage Training. We visualize samples generated by physics-constrained consistency training on 2D toy manifolds (Circle, Ellipse, Double Ellipse). (Left) Directly enforcing constraints during consistency training leads to unstable dynamics and severe mode collapse, covering only a small fraction of the target manifold. (Right) Our two-stage protocol first learns a stable, high-coverage distribution (Stage 1), then refines samples to satisfy geometric constraints (Stage 2), achieving both diversity and physical validity.

fails to capture the global topology of the manifold, learning only a partial segment of the solution space.

To address this instability, we adopt a two-stage training protocol, validated in preliminary studies [3]:

- (1) **Stage 1 (Distribution Learning):** The model is first trained to approximate the data distribution $p(\mathbf{a}, \mathbf{u})$ without explicit physical constraints. We utilize the standard adaptive sCM objective [15], which minimizes the variance of the consistency update. The loss is defined as:

$$\mathcal{L}_{\text{sCM}}(\theta, \phi) = \mathbb{E}_{\mathbf{x}_0, \mathbf{z}, t} \left[\frac{e^{w_\phi(t)}}{D} \left\| F_\theta\left(\frac{\mathbf{x}_t}{\sigma_d}, t\right) - \mathbf{y}_t \right\|_2^2 - w_\phi(t) \right], \quad (5)$$

where $w_\phi(t)$ is a learnable weighting function and $\mathbf{y}_t$ is the target estimate computed using the stop-gradient parameters θ^- . Crucially, $\mathbf{y}_t$ includes a first-order correction based on the analytic tangent $\dot{\mathbf{x}}_t = \cos(t)\mathbf{z} - \sin(t)\mathbf{x}_0$:

$$\mathbf{y}_t \triangleq F_{\theta^-}\left(\frac{\mathbf{x}_t}{\sigma_d}, t\right) + \cos(t) \left[\frac{\partial f_{\theta^-}}{\partial t} + \nabla_{\mathbf{x}} f_{\theta^-} \cdot \dot{\mathbf{x}}_t \right]. \quad (6)$$

This stage acts as a "warm-up," establishing a robust mapping for the joint distribution of coefficients and solutions before physical constraints are applied.

- (2) **Stage 2 (Physics-Informed Fine-tuning):** Once the generative backbone is established, the model is fine-tuned using the composite physics loss. In this stage, we freeze the coefficient decoder (as detailed in Sec. 3.3) and simplify the consistency term to focus on projecting generated samples onto the valid physical manifold.

Optimization Bias and Mode Collapse. While Stage 1 establishes a valid data distribution, the introduction of the physics loss in Stage 2 creates a new optimization landscape that often favors degenerate solutions. The PDE residual $\mathcal{R}(\mathbf{u}, \mathbf{a})$ is distribution-agnostic; it is minimized by *any* valid pair, including low-complexity states (e.g., constant coefficient fields) that are topologically simpler than the

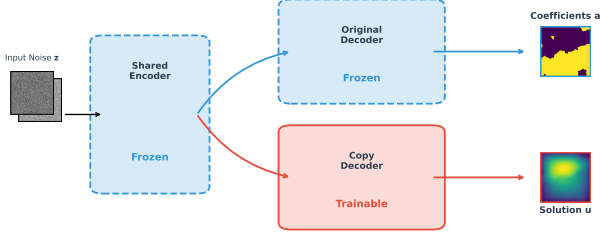


Figure 2: Architecture of the Structure-Preserving Consistency Model. To mitigate coefficient mode collapse during physics-informed fine-tuning, we employ a split-decoder architecture. The Frozen Backbone (blue, dashed) consists of the shared encoder and the Stage 1 decoder, whose parameters are fixed to preserve the learned coefficient distribution $\mathbf{a}$. The Active Branch (red, solid) is a trainable decoder that generates the solution field $\mathbf{u}$ conditioned on the fixed latent representation. The final output is the joint state $\mathbf{x} = (\mathbf{a}, \mathbf{u})$.

ground truth data. Consequently, unconstrained gradient descent tends to shift the coefficient field $\mathbf{a}$ toward these homogeneous configurations to minimize the loss rapidly, effectively disregarding the distribution learned in the warm-up phase. This collapse, which leads to the generation of trivial solutions (as demonstrated in our ablation study in Section 4.7), necessitates the structural constraints introduced below.

3.3 Structure-Preserving Channel Partitioning

To enforce the coefficient distribution learned in Stage 1, we introduce a split-decoder architecture that aligns the network’s capacity with the causal structure of the forward problem. As illustrated in Figure 2, our model modifies the standard U-Net backbone by duplicating the decoder branch while sharing a single encoder.

Architecture Design. The network processes the input $\mathbf{x}_t$ through a shared encoder and bifurcates into two parallel decoding paths:

- **The Frozen Backbone (Blue):** This path comprises the shared encoder and the *Original Decoder*. These modules are initialized with pre-trained weights from Stage 1 and are strictly frozen. This branch is responsible solely for reconstructing the coefficient field $\mathbf{a}$. By locking these weights, we ensure that the consistency mapping for the problem definition remains immutable, effectively preventing the optimizer from deviating from the valid coefficient distribution.
- **The Active Branch (Red):** We introduce a *Copy Decoder* that duplicates the architecture of the Stage 1 decoder. This branch is fully trainable and focuses exclusively on generating the solution field $\mathbf{u}$. It receives the fixed latent embeddings from the frozen encoder and minimizes the physics residuals to generate a solution $\mathbf{u}$ that is physically consistent with the coefficients.

Forward Pass. During the forward pass, the network output is partitioned into coefficient channels $F_{\theta}^{(a)}$ and solution channels $F_{\theta}^{(u)}$. We discard the solution channels from the frozen branch and the coefficient channels from the active branch. The final generated state is constructed by concatenating the valid problem parameters from the frozen path with the optimized solution from the active path:

$$F_{\theta}(\mathbf{x}_t, t) = \text{Concat} \left(\underbrace{F_{\theta_{\text{frozen}}}^{(a)}(\mathbf{x}_t, t)}_{\text{Frozen}}, \underbrace{F_{\theta_{\text{active}}}^{(u)}(\mathbf{x}_t, t)}_{\text{Trainable}} \right). \quad (7)$$

This structural constraint guarantees that the generated pair $(\mathbf{a}, \mathbf{u})$ is physically aligned, as the solution $\mathbf{u}$ is optimized specifically for the fixed coefficient topology defined by $\mathbf{a}$.

3.4 Two-Step Consistency Sampling Operator

A naive approach to physics-informed consistency training applies the PDE residual directly to single-step predictions generated from pure noise. Empirically, however, we observe that these raw one-step samples lack spatial smoothness and exhibit significantly higher PDE residuals compared to refined solutions. Because differential operators are highly sensitive to high-frequency artifacts, minimizing physics losses on these rough approximations produces noisy gradients that destabilize physics-informed consistency training.

To mitigate this, we introduce the *Two-Step Consistency Sampling Operator*, which evaluates physics constraints after a single refinement step rather than on raw one-step samples. By performing a single re-noising and denoising step, this operator projects the sample closer to the smooth solution manifold before evaluating the physics loss. This ensures that the PDE constraints are enforced on structurally valid candidates rather than noisy artifacts.

Definition 3.1 (Two-Step Consistency Sampling Operator). Let $f_{\theta}(\mathbf{x}, t)$ be the consistency model. For initial noise $\mathbf{z}$ and auxiliary noise $\mathbf{z}'$, we define the two-step operator $\hat{f}_{\theta}$ as:

$$\hat{f}_{\theta}(\mathbf{z}, \mathbf{z}'; T, t') \triangleq f_{\theta} \left(\underbrace{\sin(t')\mathbf{z}' + \cos(t')f_{\theta}(\mathbf{z}, T)}_{\text{renoised state } \mathbf{x}_{t'}}, t' \right). \quad (8)$$

This operator first produces a coarse estimate using $f_{\theta}(\mathbf{z}, T)$, then re-injects noise and applies the consistency model again at an intermediate time t' . Compared to one-step sampling, this procedure yields intermediate states that are empirically closer to the data manifold while remaining stochastic. As a result, PDE residuals evaluated on $\hat{f}_{\theta}$ exhibit reduced gradient variance, which in practice leads to more stable and effective physics-informed training.

We next incorporate this operator into a physics-informed Stage-2 objective that balances consistency preservation with robust enforcement of PDE constraints.

In contrast to approaches such as PBFM [1] that require memory-intensive multi-step ODE unrolling to mitigate discretization errors, our model natively maps to the solution manifold, bounding evaluations without unrolling bottlenecks.

3.5 Physics-Informed Stage 2 Objective

Our Stage 2 training objective fine-tunes the model to satisfy physics constraints while preserving the generative consistency learned in Stage 1. Unlike the adaptive variance minimization used in Stage 1, we employ an unweighted consistency loss to serve as a stable reference during physics optimization, avoiding the risk that adaptive weights suppress consistency precisely when physics gradients are strongest.

Unweighted Consistency Term ($\ell_{\text{sCM-L2}}$). We define the Stage 2 consistency loss as the squared Euclidean error between the model prediction and the consistency target. Consistent with our split-architecture, this loss is applied exclusively to the solution channels $\mathbf{u}$:

$$\ell_{\text{sCM-L2}}(\theta; \mathbf{x}_t, t) = \left\| F_{\theta}^{(u)}\left(\frac{\mathbf{x}_t}{\sigma_d}, t\right) - \mathbf{y}_t^{(u)} \right\|_2^2, \quad (9)$$

where $\mathbf{y}_t^{(u)}$ corresponds to the solution channels of the target estimate defined in Eq. (5). By removing the adaptive weighting terms, we enforce a uniform consistency constraint, preventing the weighting function from diminishing the loss contribution at time steps where physical constraints are most active.

Total Objective and Self-Adaptive Balancing. We integrate this consistency regularization with the physics constraints. To manage the gradient conflict between the consistency loss and the PDE residuals, we employ the Self-Adaptive weighting strategy (SA-PINN) [17] to automatically balance competing objectives.

We introduce learnable mask parameters $\lambda = \{\lambda_0, \lambda_1, \lambda_2\}$ and formulate the training as a Min-Max game. Importantly, for computational efficiency and tighter gradient coupling, the consistency term and the single-step PDE term share the same sampling trajectory:

$$\begin{aligned} \mathcal{L}_{\text{total}}(\theta, \lambda) = & \mathbb{E}_{\mathbf{x}_0, \mathbf{z}, t} \left[\underbrace{\sigma(\lambda_0) \ell_{\text{sCM-L2}}(\theta; \mathbf{x}_t, t)}_{\text{Consistency Loss}} + \underbrace{\sigma(\lambda_1) \left\| \mathcal{R}(f_{\theta}(\mathbf{x}_t, t)) \right\|_p^p}_{\text{Single-Step PDE Residual}} \right] \\ & + \mathbb{E}_{\mathbf{z}, \mathbf{z}', t'} \left[\underbrace{\sigma(\lambda_2) \left\| \mathcal{R}(\hat{f}_{\theta}(\mathbf{z}, \mathbf{z}'; T, t')) \right\|_p^p}_{\text{Two-Step Robustness}} \right], \quad (10) \end{aligned}$$

where $\sigma(\lambda) = \frac{1}{1+e^{-\lambda}}$ is the sigmoid gating function, and $\|\cdot\|_p$ denotes the L^p norm. The single-step term preserves local physical fidelity along the consistency trajectory, while the two-step term enforces robustness on denoised, structurally valid samples.

Optimization Dynamics. The training follows an adversarial update rule:

$$\theta \leftarrow \theta - \eta_{\theta} \nabla_{\theta} \mathcal{L}_{\text{total}}, \quad \lambda_i \leftarrow \lambda_i + \eta_{\lambda} \nabla_{\lambda_i} \mathcal{L}_{\text{total}}. \quad (11)$$

This mechanism allows the loss weights to dynamically adapt to the training dynamics. For instance, if the model begins to drift from the consistency manifold (increasing $\ell_{\text{sCM-L2}}$) to satisfy the physics, the gradient ascent step increases $\sigma(\lambda_0)$, forcing the model to re-prioritize consistency in subsequent updates.

Because the scalars λ_i update at the end of the computational graph, SA-PINN requires zero extra backward passes. This structural resolution bypasses the massive computational overhead of

gradient-conflict methods like CONFIG [13], which can triple memory usage and training time.

Theoretical Motivation for the Two-Step Residual. Our motivation for the two-step residual is structural, not only empirical. Let $\mathbf{x}$ be the first-step prediction, and define the refined second-step sample

$$\mathbf{Y}_{\mathbf{x}, t'} := f_{\theta}(a_{t'} \mathbf{x} + b_{t'} \mathbf{z}, t'), \quad \mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}). \quad (12)$$

where $a_{t'}$ and $b_{t'}$ are the scalar coefficients defined by the TrigFlow schedule. If we consider the L^2 penalty ($p = 2$), the conditional two-step physics loss is

$$\mathcal{L}_2(\mathbf{x}, t') := \mathbb{E}_{\mathbf{z}} [\|\mathcal{R}(\mathbf{Y}_{\mathbf{x}, t'})\|^2]. \quad (13)$$

It admits the exact decomposition

$$\mathcal{L}_2(\mathbf{x}, t') = \left\| \mathbb{E}_{\mathbf{z}} [\mathcal{R}(\mathbf{Y}_{\mathbf{x}, t'})] \right\|^2 + \text{tr}(\text{Cov}_{\mathbf{z}}(\mathcal{R}(\mathbf{Y}_{\mathbf{x}, t'}))). \quad (14)$$

Thus, unlike the one-step loss $\|\mathcal{R}(\mathbf{x})\|^2$, the two-step objective penalizes both the mean PDE violation after refinement and its variability under the model's own renoise–denoise kernel. The second term is important: a sample may have small residual at a single point yet become non-physical after a small self-consistency perturbation, which the one-step loss cannot detect.

Moreover, if f_{θ} and $\mathcal{R}$ are differentiable and $b_{t'}$ is small, then

$$\text{tr}(\text{Cov}_{\mathbf{z}}(\mathcal{R}(\mathbf{Y}_{\mathbf{x}, t'}))) = b_{t'}^2 \left\| \mathbf{J}_{\mathcal{R}}(f_{\theta}(a_{t'} \mathbf{x}, t')) \mathbf{D}_{\mathbf{x}} f_{\theta}(a_{t'} \mathbf{x}, t') \right\|_F^2 + o(b_{t'}^2). \quad (15)$$

So the two-step loss adds a first-order stability penalty on refinement directions that amplify PDE violation. Geometrically, near a regular physical point $\mathbf{y} \in \mathcal{S} := \{\mathbf{u} : \mathcal{R}(\mathbf{u}) = 0\}$, the tangent space is $T_{\mathbf{y}}\mathcal{S} = \ker \mathbf{J}_{\mathcal{R}}(\mathbf{y})$; hence this term measures how much the local renoise–denoise dynamics moves away from the tangent space of the physical solution manifold. In this sense, the two-step residual favors not only low PDE error, but also local physical stability of the refinement trajectory, which is especially desirable for few-step consistency sampling.

In principle, one could impose analogous penalties at more intermediate steps, but this would substantially increase training cost. Empirically, we found that the 2-step loss gives a good tradeoff between efficiency and performance.

3.6 Inference and Forward Problem Solving

At inference time, our framework supports two distinct modes of operation. When sampling random solution–coefficient pairs from the learned joint distribution, we directly apply the standard consistency model sampling procedure, enabling fast generation via single-step denoising updates.

In contrast, solving the forward problem—where the coefficient field $\mathbf{a}$ is given and the solution $\mathbf{u}$ is sought—requires a conditional generation mechanism. In this setting, naively applying unconditional sampling is insufficient, as the model must respect the prescribed coefficients while satisfying the governing PDE. We therefore reformulate the forward problem as a *zero-shot inpainting* task.

We introduce an iterative inference procedure that treats the consistency model as a learned solver. Leveraging the zero-shot inpainting strategy proposed in the original Consistency Models framework [21], we strictly enforce the known coefficient channels

$\mathbf{a}_{\text{obs}}$ via a hard projection operator $\mathcal{P}$ at each denoising step. This guides the generative trajectory to evolve on the manifold defined by the specific physical instance. Formally, for a generated state $\hat{\mathbf{x}}$ and a binary mask $\mathbf{M}$ (where $\mathbf{M} = 1$ for coefficient channels and 0 otherwise), the projection is defined as:

$$\mathcal{P}(\hat{\mathbf{x}}, \mathbf{x}_{\text{obs}}) = \mathbf{M} \odot \mathbf{x}_{\text{obs}} + (\mathbf{1} - \mathbf{M}) \odot \hat{\mathbf{x}}. \quad (16)$$

Importantly, this procedure does not involve explicit numerical optimization or gradient-based PDE residual evaluation at inference time. Instead, the trained model implicitly encodes the coupling between $\mathbf{a}$ and $\mathbf{u}$. The iterative "predict-project-renoise" updates act to align the generated solution with the imposed coefficients, leveraging the robustness induced by our two-step training objective to correct approximation errors.

Algorithmic details of this projection-based forward solver are summarized in Appendix A.

4 Experiments

We evaluate the efficacy of our framework on several canonical steady-state PDE benchmarks (Darcy flow, Poisson, and Helmholtz equations), as well as the time-dependent Navier-Stokes equations. We adopt the pre-computed datasets released by the DiffusionPDE authors [6]. These datasets were generated using high-fidelity finite element method (FEM) solvers on a 128×128 spatial grids, providing 50,000 paired samples for each equation family. Unless otherwise specified, we utilize the L^2 norm ($p = 2$) for the PDE residual penalty across all evaluations.

4.1 PDE Problem Settings

We consider a suite of standard elliptic PDEs alongside the time-dependent Navier-Stokes equations. Full mathematical definitions are provided in Appendix B.

4.2 Evaluation Metric: Relative H^1 Error

For forward solution prediction in elliptic PDE benchmarks, prior work often relies solely on the relative L^2 error to quantify solution accuracy. However, this metric can be insufficient for physics-informed assessment, as it measures pointwise proximity but is insensitive to high-frequency oscillations or gradient mismatches. As a result, a candidate solution may achieve a low L^2 error while exhibiting non-physical roughness that violates the governing differential equation.

To more rigorously assess physical validity in forward solution prediction, we adopt the relative H^1 error. This metric accounts for both solution values and their spatial derivatives, ensuring that predicted fields are not only accurate but also smooth and consistent with the underlying physics. Different metrics are employed for conditional inference and unconditional sampling, and are introduced in the corresponding experimental sections.

Definition. For a predicted solution $\hat{u}$ and ground-truth solution u defined on a domain Ω , the squared H^1 norm is defined as

$$\|u\|_{H^1(\Omega)}^2 \triangleq \|u\|_{L^2(\Omega)}^2 + \|\nabla u\|_{L^2(\Omega)}^2 = \int_{\Omega} |u|^2 dx + \int_{\Omega} |\nabla u|^2 dx. \quad (17)$$

The relative H^1 error is then given by

$$\mathcal{E}_{H^1} = \frac{\|\hat{u} - u\|_{H^1(\Omega)}}{\|u\|_{H^1(\Omega)}}. \quad (18)$$

Implementation. Since the solutions are defined on a discrete grid with spacing h , spatial derivatives are approximated using second-order central differences at interior points and first-order differences at the boundaries. The discrete H^1 norm used in our evaluation is

$$\|e\|_{H^1}^2 \approx \sum_{i,j} (|e_{i,j}|^2 + |\delta_x e_{i,j}|^2 + |\delta_y e_{i,j}|^2) h^2, \quad (19)$$

where $e = \hat{u} - u$ denotes the error field, and δ_x, δ_y are discrete central difference operators (e.g., $\delta_x e_{i,j} \approx \frac{e_{i+1,j} - e_{i-1,j}}{2h}$).

4.3 Implementation and Baselines

To assess the impact of our physics-informed fine-tuning, we compare the proposed method against its pre-trained base model and established diffusion-based solvers.

Network Architecture. To ensure a fair comparison with the baseline, we adopt the U-Net backbone used in DiffusionPDE [6], which is based on the DDPM++ architecture [23]. We modify only the parameterization and normalization components relevant to consistency training. Specifically, while DiffusionPDE employs the EDM formulation [9] with diffusion-specific pre-conditioning, we remove these scaling factors and operate directly on the underlying U-Net trunk. We then apply the TrigFlow parameterization described in Sec. 3.1.

To improve training stability, we further adjust the normalization layers. As observed in the sCM framework [15], the standard AdaGN layer can lead to divergence in consistency models; accordingly, we replace it with a modified Adaptive Double Normalization that incorporates pixel normalization [8].

Training Protocol. Training is conducted in three sequential phases. We employ the AdamW optimizer [14] throughout all stages, with hyperparameters summarized in Appendix C.

- (1) **Diffusion Pre-training:** We first train a score-based diffusion model to approximate the data distribution $p_{\text{data}}(\mathbf{a}, \mathbf{u})$.
- (2) **Stage 1 (Consistency Training):** We initialize the consistency model from the pre-trained diffusion weights and optimize it using the standard sCM objective (Eq. 5) to enforce self-consistency.
- (3) **Stage 2 (Physics Fine-Tuning):** Starting from the Stage 1 checkpoint, we fine-tune the model using the physics-informed joint Min-Max objective (Eq. 10), with the coefficient decoder frozen. The balancing parameters λ are initialized to λ_{init} (Table 4) and optimized jointly with the model weights.

Metric: Number of Function Evaluations (NFE). To quantify inference efficiency, we report the Number of Function Evaluations (NFE) rather than raw sampling steps. This metric directly reflects the number of neural network forward passes and is independent of the specific numerical solver. DiffusionPDE employs Heun's second-order solver [9] together with Diffusion Posterior Sampling (DPS) [4]. This predictor-corrector scheme requires two network evaluations per sampling step, leading to an NFE of $2T - 1$ for T

sampling steps. In contrast, our sCM-PINN supports single-pass generation, requiring only one network evaluation per step ($\text{NFE} = T$).

4.4 Forward Problem Results

Table 1 summarizes the quantitative results for forward PDE solving. All reported metrics are computed as averages over 256 randomly selected test samples. We evaluate the models under varying computational budgets to examine the trade-off between inference efficiency and physical accuracy.

Accuracy and Stability. Our physics-informed model (sCM-PINN) consistently outperforms the Stage 1 data-only baseline (sCM) across all three physical systems. For the Darcy Flow benchmark, physics-informed fine-tuning reduces the relative H^1 error by approximately 42% at 65 steps (1.05×10^{-1} vs. 1.81×10^{-1}), demonstrating that the frozen-decoder strategy successfully injects physical constraints without degrading the learned data distribution. Importantly, sCM-PINN also achieves accuracy comparable to or better than the DiffusionPDE baseline while operating at substantially fewer function evaluations.

Computational Efficiency. Thanks to the consistency formulation, sCM-PINN converges rapidly to physically accurate solutions. On the Darcy Flow benchmark, sCM-PINN reaches a relative error of 1.10×10^{-1} in only 17 evaluations, whereas DiffusionPDE exhibits a substantially higher error (3.49×10^{-1}) even after 127 evaluations. This corresponds to an effective inference speedup of nearly $7.5\times$ in terms of network evaluations. On the Poisson equation, sCM-PINN attains a relative error of 1.81×10^{-1} with 65 evaluations, slightly improving upon DiffusionPDE with 127 evaluations (1.85×10^{-1}).

Table 1: Comparison of relative H^1 error for the forward problem averaged over 256 test samples. The best results are highlighted in bold. NFE denotes Number of Function Evaluations.

Method	NFE	Darcy Flow	Poisson	Helmholtz	Navier-Stokes
DiffusionPDE	31	6.37×10^{-1}	9.99×10^{-1}	1.38×10^0	1.13×10^0
	63	4.18×10^{-1}	5.34×10^{-1}	1.20×10^0	8.65×10^{-1}
	127	3.49×10^{-1}	1.85×10^{-1}	9.37×10^{-1}	3.85×10^{-1}
sCM	17	2.00×10^{-1}	4.17×10^{-1}	4.31×10^{-1}	2.79×10^{-1}
	33	1.92×10^{-1}	3.33×10^{-1}	3.34×10^{-1}	2.40×10^{-1}
	65	1.81×10^{-1}	2.78×10^{-1}	2.82×10^{-1}	2.12×10^{-1}
sCM-PINN (Ours)	17	1.10×10^{-1}	3.20×10^{-1}	3.76×10^{-1}	2.06×10^{-1}
	33	1.08×10^{-1}	2.28×10^{-1}	2.87×10^{-1}	1.62×10^{-1}
	65	1.05×10^{-1}	1.81×10^{-1}	2.33×10^{-1}	1.43×10^{-1}

Wall-Clock Efficiency. While the number of function evaluations (NFE) provides a hardware-agnostic measure of sampling complexity, it does not fully capture the computational cost of diffusion-based solvers that rely on gradient-based guidance. In particular, DiffusionPDE employs diffusion posterior sampling (DPS) [4], where each sampling step incurs the additional overhead of backpropagating through the PDE residual operator.

To account for this discrepancy, we additionally report wall-clock inference time (See Appendix D for details), comparing sCM-PINN with DiffusionPDE for generating a single sample (batch size = 1). Despite using a similar number of steps (65 vs. 63), sCM-PINN

achieves substantially lower inference time (1.56s vs. 3.86s), as each step consists solely of a forward network evaluation, whereas DPS requires repeated gradient computations. Furthermore, we benchmarked sCM-PINN against a traditional CPU Immersed Interface Method (IIM) solver [11] for Darcy flow. Averaged over 100 samples, the IIM solver takes 1.40s per sample, demonstrating that sCM-PINN reaches comparable practical clock time to a traditional numerical solver, while additionally supporting joint generative modeling.

4.5 Conditional Source Reconstruction Results

We evaluate the framework’s ability to perform *conditional generative inference*, in which the source term $a(\mathbf{x})$ is inferred from partial or complete observations of the solution field $u(\mathbf{x})$. This task corresponds to an ill-posed inverse setting, in which multiple source configurations may be consistent with the same observed solution. We consider Poisson and Helmholtz benchmarks and report quantitative results in Table 2.

Metric Selection (Relative L^2 vs. H^1). Unlike the forward problem, we adopt the relative L^2 error as the primary metric for conditional inference. This choice reflects the ill-posed nature of inverse PDE settings, which are inherently sensitive to high-frequency perturbations and may admit multiple valid reconstructions. Moreover, while the solution field u is typically smooth due to elliptic regularity, the source term a is not necessarily differentiable and may contain sharp transitions or discontinuities.

Formally, the H^1 norm requires a function and its first weak derivative to be both square-integrable (i.e., both $u \in L^2$ and $\nabla u \in L^2$). Specifically, for the solution u of the Poisson and Helmholtz equations to exist in H^1 , the source term a only needs to belong to the dual space H^{-1} . Even if $a \in L^2$ (which implies $u \in H^2$), the source term a might not belong to H^1 .

As a result, an H^1 -based metric would disproportionately penalize high-frequency components and overemphasize small-scale artifacts, leading to a misleading assessment of reconstruction quality. We emphasize that our goal is not to recover a unique ground-truth source, but rather to generate physically consistent source realizations conditioned on the observed solution.

Table 2: Comparison of relative L^2 error for conditional source inference averaged over 256 test samples. The best results are highlighted in bold. NFE denotes the Number of Function Evaluations.

Method	NFE	Poisson	Helmholtz
DiffusionPDE	31	1.21×10^0	1.34×10^0
	63	8.91×10^{-1}	1.20×10^0
	127	4.56×10^{-1}	9.74×10^{-1}
sCM-PINN (Ours)	17	5.92×10^{-1}	5.72×10^{-1}
	33	4.85×10^{-1}	4.58×10^{-1}
	65	3.99×10^{-1}	3.81×10^{-1}

4.6 Unconditional Sampling

Beyond forward problem solving and conditional inference, we evaluate the quality of *unconditional* samples generated by each model. In this setting, samples are drawn directly from the learned joint distribution $p(\mathbf{a}, \mathbf{u})$ without conditioning on observations. This experiment assesses whether the generative prior learned by each method captures both the data distribution and the governing physical constraints in a free-form generation regime.

For unconditional sampling, the exact solution is not available, so evaluation must rely on PDE residuals rather than solution errors. In this setting, an H^1 -type error is not well-defined, and it is more appropriate to measure the L^2 norm of the residual to assess physical consistency.

To quantify physical fidelity, we compute the normalized PDE residual $\|\mathcal{R}\|_2 \cdot h^2$, where $\mathcal{R}$ denotes the residual operator and h is the grid spacing. Lower values indicate closer adherence to the governing equations. Since unconditional samples are not constrained by observations, this metric evaluates physical plausibility rather than exact solution accuracy, reflecting the quality of the learned generative prior.

Table 3 presents the quantitative comparison. Notably, sCM-PINN achieves the lowest residual errors across all three benchmarks while requiring only 2 function evaluations (NFE). This suggests that physics-informed fine-tuning improves not only conditional accuracy but also aligns the learned generative prior more closely with the underlying physical laws. In contrast, the standard sCM baseline exhibits substantially higher residuals (up to 20× higher on Darcy Flow), indicating that data-driven training alone struggles to enforce physical consistency in unconditional generation. While DiffusionPDE performs robustly, sCM-PINN matches or exceeds its physical fidelity with a 30× reduction in computational cost.

We visualize unconditional samples for Darcy Flow in Figure 3. Qualitatively, sCM-PINN produces sharp, coherent coefficient interfaces and smooth solution fields that closely resemble the reference data. The residual maps further reveal that DiffusionPDE tends to produce localized high-error regions in high-permeability zones ($a = 12$), whereas sCM-PINN maintains more uniform physical consistency across the domain.

Table 3: Comparison of unconditional sampling quality. We report the normalized PDE residual ($\|\mathcal{R}\|_2 \cdot h^2$), where lower values indicate better physical fidelity.

Method	NFE	Darcy Flow	Poisson	Helmholtz	Navier-Stokes
DiffusionPDE	63	3.22×10^{-2}	1.18×10^{-2}	1.99×10^{-2}	1.92×10^{-2}
sCM	2	4.44×10^{-1}	4.05×10^{-2}	4.67×10^{-2}	2.12×10^{-2}
sCM-PINN (Ours)	2	2.00×10^{-2}	1.13×10^{-2}	1.11×10^{-2}	7.52×10^{-3}

4.7 Ablation Study: Effect of Freezing the Coefficient Decoder

In this section, we investigate the role of freezing the coefficient decoder during Stage 2 physics-informed fine-tuning. We compare our proposed approach against a baseline where the entire model is trained jointly without freezing any components. For this

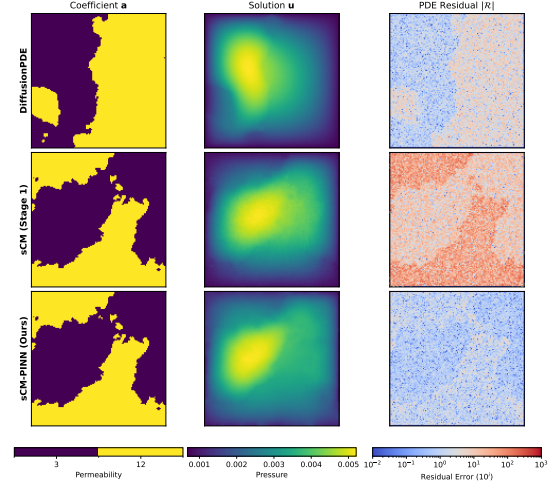


Figure 3: Unconditional Sampling Quality (Darcy Flow). Visual comparison of generated coefficient fields $\mathbf{a}$ (left), solution fields $\mathbf{u}$ (middle), and absolute PDE residuals $|\mathcal{R}|$ (right). The baseline sCM (middle row) exhibits visible artifacts, whereas sCM-PINN (bottom row) produces sharper coefficient interfaces and solution fields with lower and more uniform residuals, comparable to those obtained by the 63-step DiffusionPDE baseline (top row) while requiring only 2 function evaluations. Note the reduced residual variance in the sCM-PINN samples relative to the diffusion baseline.

joint training baseline, we modify the Simplified Consistency Term (Eq. 9) in the total loss to explicitly supervise the coefficient channel. Specifically, we replace the component-wise masked loss with a full-state L_2 loss over both coefficient and solution channels:

$$\ell_{\text{sCM-L2-joint}}(\theta; \mathbf{x}_t, t) = \left\| F_{\theta} \left(\frac{\mathbf{x}_t}{\sigma_d}, t \right) - \mathbf{y}_t \right\|_2^2.$$

Coefficient Mode Collapse. In the Darcy Flow problem, the permeability field $\mathbf{a}(\mathbf{x})$ takes discrete values in $\{3, 12\}$ with approximately equal probability. When we allow the coefficient decoder to update during physics fine-tuning, the model rapidly collapses to a unimodal distribution concentrated at $a \approx 12$. We attribute this behavior to the observation that higher permeability values yield smoother pressure fields $\mathbf{u}$ for a fixed forcing term. Consequently, the optimizer minimizes the PDE residual by shifting the coefficient distribution toward these simpler configurations, effectively ignoring the target data distribution learned in Stage 1.

Quantitative Analysis. Figure 4 compares the pixel-wise histograms and spatial coefficient maps of the generated fields. We generate 128 samples from each model using two-step sampling, visualizing both pixel-wise value distributions and representative spatial coefficient maps. The joint training baseline (red) exhibits severe mode collapse, failing to capture the low-permeability mode properly. In contrast, our frozen decoder approach (blue) preserves the ground-truth statistics (gray) while still minimizing the PDE residual. These results indicate that disentangling input generation (Stage 1, fixed) from solution optimization (Stage 2, active) is critical

for preserving the learned coefficient distribution while enforcing physical consistency.

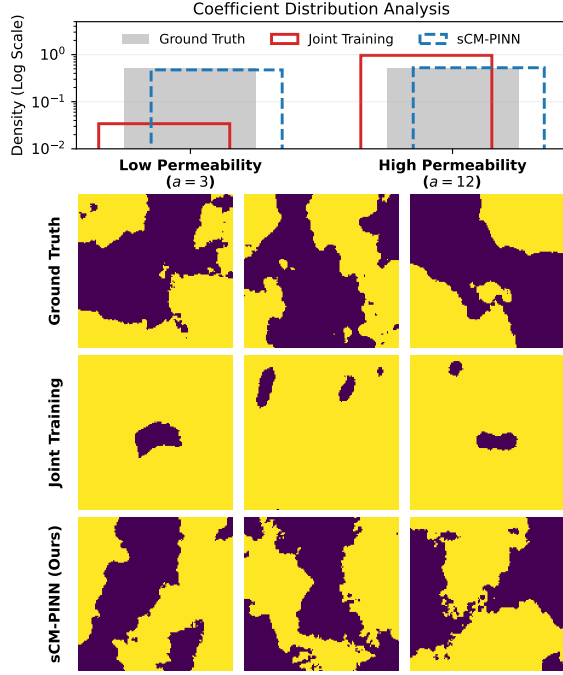


Figure 4: Analysis of Coefficient Mode Collapse. (Top) Pixel-wise histogram showing the bimodal ground-truth distribution (gray). Joint training (red) collapses the distribution to a single mode at $a \approx 12$. (Bottom) Representative spatial samples. The joint-training baseline (middle row) produces only high-permeability features, whereas the frozen decoder strategy (bottom row) recovers the full coefficient structure.

4.8 Extended Experimental Results

To broaden the scope of our evaluation, we conducted extended experiments (Appendix E).

Deterministic Operators: We compare against FNO [12] and DeepONet [16] for forward problems. While FNO excels on Poisson and Helmholtz, it is a strictly supervised model. Conversely, sCM-PINN natively supports stochastic modeling and outperforms FNO on Darcy Flow (Table 6).

Distributional Fidelity: We benchmark against concurrent physics-informed models (PBFM, PIDM) on their specific Darcy Flow dataset. Evaluated via Wasserstein distance (WD) and Jensen-Shannon divergence (JS), sCM-PINN maintains distributional quality orders of magnitude better than these baselines at just 2 NFE (Table 7).

Navier-Stokes System: Building upon the forward and unconditional results presented earlier, we provide extended evaluations for this time-dependent system in Appendix E.3. For the inverse problem (Table 8), sCM-PINN achieves a lower H^1 error than DiffusionPDE in fewer steps. For unconditional generation (Table 9), we analyze the computational trade-off, showing sCM-PINN achieves

the lowest PDE residual at just 2 NFE, and recovers highly competitive distributional fidelity (WD and JS) when scaled to 17 NFE.

Sensitivity and Training Cost: Our method transfers robustly without exhaustive tuning. Transitioning to Stage 2 at varying checkpoints yields stable metrics, proving precise scheduling is unnecessary (Table 10). Total training requires ~ 17 hours on an RTX 4090.

5 Conclusion

In this work, we introduced *sCM-PINN*, a physics-informed consistency modeling framework for efficient PDE solving and physically grounded generative modeling. We identified key stability challenges that arise when enforcing physical constraints in consistency training, where naive joint optimization can lead to degenerate solutions and coefficient mode collapse. To address this issue, we proposed a structure-preserving training strategy based on a two-stage protocol that couples a frozen coefficient decoder with a physics-informed residual objective. This design stabilizes optimization, enabling the model to respect physical laws while preserving the learned generative prior.

Across Darcy Flow, Poisson, and Helmholtz benchmarks, sCM-PINN consistently achieves strong physical accuracy with substantially fewer function evaluations than diffusion-based PDE solvers. In particular, sCM-PINN attains the lowest relative H^1 error among the evaluated methods while requiring only a small number of model evaluations at inference. Ablation studies further demonstrate that freezing the coefficient decoder during physics fine-tuning is critical for preventing mode collapse and maintaining the diversity of the learned coefficient distribution. Moreover, unconditional sampling experiments show that sCM-PINN produces physically plausible samples with only two function evaluations, whereas diffusion-based baselines require orders of magnitude more computational effort. By bridging the efficiency of consistency models with the rigor of physics-informed learning, sCM-PINN offers a promising direction toward real-time scientific simulation and generative modeling of physical systems. Future work will explore extensions to more complex PDEs, time-dependent dynamics, and higher-dimensional settings.

6 Limitations and Ethical Considerations

Our multi-phase training pipeline introduces additional implementation overhead, and the continuous-time formulation requires Jacobian-Vector Products (JVPs), increasing memory cost. As our experiments use only synthetic open-source PDE data, we foresee no ethical concerns.

Acknowledgments

T.-S. Lin acknowledges the supports by National Science and Technology Council, Taiwan, under research grants 111-2628-M-A49-008-MY4 and 114-2124-M-390-001. M.-C. Lai acknowledges the support by National Science and Technology Council, Taiwan, under research grant 113-2115-M-A49-014-MY3.

GenAI Disclosure

The authors acknowledge the use of large language models to polish and edit the written text of the manuscript.

References

- [1] Giacomo Baldan, Qiang Liu, Alberto Guardone, and Nils Thuerey. 2026. Physics vs Distributions: Pareto Optimal Flow Matching with Physics Constraints. In *International Conference on Learning Representations*.
- [2] Jan-Hendrik Bastek, WaiChing Sun, and Dennis Kochmann. 2025. Physics-Informed Diffusion Models. In *International Conference on Learning Representations*.
- [3] Che-Chia Chang, Chen-Yang Dai, Te-Sheng Lin, Ming-Chih Lai, and Chieh-Hsin Lai. 2025. Consistency Training with Physical Constraints. In *ICLR Workshop on Frontiers in Probabilistic Inference: Learning meets Sampling*.
- [4] Hyungjin Chung, Jeongsol Kim, Michael Thompson McCann, Marc Louis Klasky, and Jong Chul Ye. 2023. Diffusion Posterior Sampling for General Noisy Inverse Problems. In *International Conference on Learning Representations*.
- [5] Jonathan Ho, Ajay Jain, and Pieter Abbeel. 2020. Denoising Diffusion Probabilistic Models. In *Advances in Neural Information Processing Systems*.
- [6] Jiahe Huang, Guandao Yang, Zichen Wang, and Jeong Joon Park. 2024. Diffusion-PDE: Generative PDE-Solving under Partial Observation. In *Advances in Neural Information Processing Systems*.
- [7] Christian Jacobsen, Yilin Zhuang, and Karthik Duraisamy. 2025. CoCoGen: Physically Consistent and Conditioned Score-Based Generative Models for Forward and Inverse Problems. *SIAM Journal on Scientific Computing* 47 (2025), C399–C425. doi:10.1137/24M1636071
- [8] Tero Karras, Timo Aila, Samuli Laine, and Jaakko Lehtinen. 2018. Progressive Growing of GANs for Improved Quality, Stability, and Variation. In *International Conference on Learning Representations*.
- [9] Tero Karras, Miika Aittala, Timo Aila, and Samuli Laine. 2022. Elucidating the Design Space of Diffusion-Based Generative Models. In *Advances in Neural Information Processing Systems*.
- [10] Chieh-Hsin Lai, Yang Song, Dongjun Kim, Yuki Mitsufuji, and Stefano Ermon. 2025. The principles of diffusion models. *arXiv preprint arXiv:2510.21890* (2025).
- [11] Zhilin Li and Kazufumi Ito. 2006. *The Immersed Interface Method*. Society for Industrial and Applied Mathematics. doi:10.1137/1.9780898717464
- [12] Zongyi Li, Nikola Borislavov Kovachki, Kamyar Azizzadenesheli, Burigede Liu, Kaushik Bhattacharya, Andrew Stuart, and Anima Anandkumar. 2021. Fourier Neural Operator for Parametric Partial Differential Equations. In *International Conference on Learning Representations*.
- [13] Qiang Liu, Mengyu Chu, and Nils Thuerey. 2025. ConFIG: Towards Conflict-free Training of Physics Informed Neural Networks. In *International Conference on Learning Representations*.
- [14] Ilya Loshchilov and Frank Hutter. 2019. Decoupled Weight Decay Regularization. In *International Conference on Learning Representations*.
- [15] Cheng Lu and Yang Song. 2025. Simplifying, Stabilizing and Scaling Continuous-time Consistency Models. In *International Conference on Learning Representations*.
- [16] Lu Lu, Pengzhan Jin, Guofei Pang, Zhongqiang Zhang, and George Em Karniadakis. 2021. Learning nonlinear operators via DeepONet based on the universal approximation theorem of operators. *Nature Machine Intelligence* 3 (2021), 218–229. doi:10.1038/s42256-021-00302-5
- [17] Levi D. McClenny and Ulisses M. Braga-Neto. 2023. Self-adaptive physics-informed neural networks. *J. Comput. Phys.* 474 (2023), 111722. doi:10.1016/j.jcp.2022.111722
- [18] Maziar Raissi, Paris Perdikaris, and George E Karniadakis. 2019. Physics-informed neural networks: A deep learning framework for solving forward and inverse problems involving nonlinear partial differential equations. *J. Comput. Phys.* 378 (2019), 686–707. doi:10.1016/j.jcp.2018.10.045
- [19] Dule Shu, Zijie Li, and Amir Barati Farimani. 2023. A physics-informed diffusion model for high-fidelity flow field reconstruction. *J. Comput. Phys.* 478 (2023), 111972. doi:10.1016/j.jcp.2023.111972
- [20] Jascha Sohl-Dickstein, Eric Weiss, Niru Maheswaranathan, and Surya Ganguli. 2015. Deep Unsupervised Learning using Nonequilibrium Thermodynamics. In *International Conference on Machine Learning*.
- [21] Yang Song, Prafulla Dhariwal, Mark Chen, and Ilya Sutskever. 2023. Consistency Models. In *International Conference on Machine Learning*.
- [22] Yang Song and Stefano Ermon. 2019. Generative Modeling by Estimating Gradients of the Data Distribution. In *Advances in Neural Information Processing Systems*.
- [23] Yang Song, Jascha Sohl-Dickstein, Diederik P Kingma, Abhishek Kumar, Stefano Ermon, and Ben Poole. 2021. Score-Based Generative Modeling through Stochastic Differential Equations. In *International Conference on Learning Representations*.

A Algorithm Details

We provide the complete pseudocode for the projection-based forward solver detailed in Algorithm 1.

Algorithm 1 Measurement-Constrained Consistency Sampling (TrigFlow)

Require: Consistency Model $f_\theta(\mathbf{x}, t)$, Observation $\mathbf{x}_{\text{obs}}$, Mask $\mathbf{M}$.
Require: Time schedule $t_0 > t_1 > \dots > t_N$ (where $t_0 \approx \pi/2$, $t_N \approx 0$).
Require: Data deviation σ_d .

- 1: **Initialize:** Sample noise $\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$.
- 2: Set initial noisy state $\mathbf{x}_{t_0} \leftarrow \sigma_d \mathbf{z}$ (assuming $t_0 \approx \pi/2$).
- 3: **Initial Estimation:**
- 4: $\hat{\mathbf{x}} \leftarrow f_\theta(\mathbf{x}_{t_0}, t_0)$ {Predict $\mathbf{x}_0$ using Eq. 4}
- 5: $\hat{\mathbf{x}} \leftarrow \mathbf{M} \odot \mathbf{x}_{\text{obs}} + (\mathbf{1} - \mathbf{M}) \odot \hat{\mathbf{x}}$ {Apply measurement constraint}
- 6: **for** $n = 1$ **to** N **do**
- 7: Sample fresh noise $\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$.
- 8: **Re-noising (TrigFlow):**
- 9: $\mathbf{x}_{t_n} \leftarrow \cos(t_n)\hat{\mathbf{x}} + \sin(t_n)\sigma_d \mathbf{z}$ {Back-project to manifold at t_n }
- 10: **Consistency Update:**
- 11: $\hat{\mathbf{x}} \leftarrow f_\theta(\mathbf{x}_{t_n}, t_n)$ {Jump to data origin}
- 12: **Projection:**
- 13: $\hat{\mathbf{x}} \leftarrow \mathbf{M} \odot \mathbf{x}_{\text{obs}} + (\mathbf{1} - \mathbf{M}) \odot \hat{\mathbf{x}}$ {Enforce measurements}
- 14: **end for**
- 15: **return** $\hat{\mathbf{x}}$

B PDE Problem Details

Darcy Flow. Darcy flow describes fluid transport through porous media and serves as a standard benchmark in hydrogeology and reservoir modeling. We consider the steady-state 2D Darcy flow equation on the unit square $\Omega = (0, 1)^2$:

$$-\nabla \cdot (a(\mathbf{x}) \nabla u(\mathbf{x})) = f(\mathbf{x}), \quad \mathbf{x} \in \Omega, \quad (20)$$

where $u(\mathbf{x})$ denotes the fluid pressure and $a(\mathbf{x})$ is the permeability field. We evaluate on two commonly used Darcy-flow benchmarks that differ in their coefficient fields, forcing terms, and boundary conditions.

DiffusionPDE Benchmark. The DiffusionPDE benchmark is discretized on a 128×128 spatial grid over the unit square domain $\Omega = (0, 1)^2$, with homogeneous Dirichlet boundary conditions,

$$u(\mathbf{x}) = 0, \quad \mathbf{x} \in \partial\Omega. \quad (21)$$

The source term is fixed as $f(\mathbf{x}) = 1$. The permeability field is a piecewise-constant medium taking values in $\{3, 12\}$, representing a composite material with distinct permeability regions.

PBFM/PIDM Benchmark. We also consider a second Darcy flow benchmark adopted in recent studies on physics-informed generative modeling [1, 2]. This benchmark is discretized on a 64×64 spatial grid and imposes homogeneous Neumann boundary conditions together with a mean-zero pressure constraint:

$$\begin{aligned} \nabla u(\mathbf{x}) \cdot \hat{\mathbf{n}}(\mathbf{x}) &= 0, & \mathbf{x} \in \partial\Omega, \\ \int_{\Omega} u(\mathbf{x}) d\mathbf{x} &= 0. \end{aligned} \quad (22)$$

The permeability field is modeled as a continuous log-normal Gaussian random field, $a(\mathbf{x}) = \exp(G(\mathbf{x}))$, where $G(\mathbf{x})$ is a Gaussian

random field. The source term is given by

$$f(\mathbf{x}) = \begin{cases} r, & \text{if } |x_i - \frac{w}{2}| \leq \frac{w}{2}, \\ -r, & \text{if } |x_i - 1 + \frac{w}{2}| \leq \frac{w}{2}, \\ 0, & \text{otherwise,} \end{cases} \quad (23)$$

with $r = 10$ and $w = 0.125$.

Poisson Equation. As a fundamental elliptic operator, the Poisson equation models potential fields arising in electrostatics, gravitation, and steady-state heat conduction. We solve

$$\begin{aligned} \nabla^2 u(\mathbf{x}) &= a(\mathbf{x}), & \mathbf{x} \in \Omega, \\ u(\mathbf{x}) &= 0, & \mathbf{x} \in \partial\Omega. \end{aligned} \quad (24)$$

In this setting, $a(\mathbf{x})$ acts as a spatially varying source term, while $u(\mathbf{x})$ represents the resulting potential field. The source terms are synthesized using Gaussian random fields (GRFs), yielding a diverse collection of smooth input structures.

Inhomogeneous Helmholtz Equation. To assess the model’s ability to handle wave-like phenomena, we consider the inhomogeneous Helmholtz equation, which augments the Poisson operator with a reaction term:

$$\begin{aligned} \nabla^2 u(\mathbf{x}) + k^2 u(\mathbf{x}) &= a(\mathbf{x}), & \mathbf{x} \in \Omega, \\ u(\mathbf{x}) &= 0, & \mathbf{x} \in \partial\Omega. \end{aligned} \quad (25)$$

We fix the wavenumber at $k = 1$. As in the Poisson case, $a(\mathbf{x})$ denotes the source distribution. The addition of the linear reaction term $k^2 u(\mathbf{x})$ introduces oscillatory behavior in the solution manifold, posing a qualitatively different challenge from the purely diffusive Darcy and Poisson systems.

Non-bounded Navier-Stokes Equation. We consider the incompressible Navier-Stokes equation for the vorticity.

$$\begin{aligned} \partial_t w(\mathbf{x}, t) + v(\mathbf{x}, t) \cdot \nabla w(\mathbf{x}, t) &= \nu \Delta w(\mathbf{x}, t) + q(\mathbf{x}), & \mathbf{x} \in \Omega, t \in (0, T] \\ \nabla \cdot v(\mathbf{x}, t) &= 0, & \mathbf{x} \in \Omega, t \in [0, T] \end{aligned} \quad (26)$$

Here $w = \nabla \times v$ is the vorticity, $v(\mathbf{x}, t)$ is the velocity at $\mathbf{x}$ at time t , and $q(\mathbf{x})$ is a force field. The viscosity coefficient is set to $\nu = 10^{-3}$, which corresponds to the Reynolds number $Re = \frac{1}{\nu} = 1000$.

C Training Hyperparameter Settings

We provide the training hyperparameters for our training protocol, which is summarized in Table 4.

Table 4: Training hyperparameters. LR: Learning Rate, BS: Batch Size. Epochs are listed for (Darcy / Poisson / Helmholtz / Navier-Stokes). λ_{init} denotes the initialization of the learnable loss masks $\{\lambda_0, \lambda_1, \lambda_2\}$ in Stage 2.

Stage	LR	BS	λ_{init}	Epochs (D/P/H/N)
Pre-train	1e-3	16	–	10 / 10 / 10 / 10
Stage 1	1e-3	2	–	8 / 8 / 8 / 8
Stage 2	1e-4	2	$\{10, -10, -10\}$	1 / 2 / 2 / 2

D Hardware and Runtime Analysis

All inference experiments were conducted on a workstation equipped with an Intel i7-13700 CPU and an NVIDIA RTX 4090 GPU with 24GB of VRAM.

Table 5 compares the wall-clock time required to generate a single Darcy flow solution. First, we compare against the generative baseline. sCM-PINN significantly outperforms DiffusionPDE (1.56s vs. 3.86s) by eliminating the computational overhead of the iterative gradient-based guidance used in DiffusionPDE.

Furthermore, to demonstrate practical utility, we benchmarked our sCM-PINN against a traditional CPU Immersed Interface Method (IIM) solver [11] for Darcy flow. Averaged over 100 samples, the IIM solver takes 1.40 seconds per sample. At our maximum forward-problem setting (65 steps, batch size 1), sCM-PINN achieves comparable wall-clock time at 1.56 seconds. We note that since the IIM runs on the CPU and sCM-PINN runs on the GPU, the precise timing ratio will vary with different hardware configurations.

Table 5: Wall-clock sampling time (in seconds) for generating a single Darcy flow sample.

Method	Hardware	NFE	Time (s)
IIM (Traditional Solver)	CPU (i7-13700)	–	1.40
sCM-PINN (Ours)	GPU (RTX 4090)	65	1.56
DiffusionPDE	GPU (RTX 4090)	63	3.86

E Extended Experimental Results

This section provides the comprehensive quantitative tables supporting the additional evaluations discussed in Sec. 4.8.

E.1 Comparison to Deterministic Neural Operators

To show that our generative approach is competitive in accuracy for forward problems, we compared sCM-PINN against deterministic neural operator baselines, specifically FNO and DeepONet. As shown in Table 6, sCM-PINN (at NFE=65) achieves competitive forward-solving relative L^2 errors against deterministic baselines and DiffusionPDE, outperforming FNO on the Darcy Flow benchmark.

While FNO excels on Poisson/Helmholtz, it is a strictly supervised, deterministic solver tailored solely for forward tasks. This deterministic nature makes FNO more restrictive for uncertainty quantification, inverse problems, and learning from partial or scarce observations. Conversely, sCM-PINN is a unified generative model that natively supports stochastic modeling, unconditional sampling, forward solving, and source reconstruction.

Table 6: Relative L^2 error comparison against deterministic neural operators.

Method	Darcy Flow	Poisson	Helmholtz
DeepONet	1.23×10^{-1}	1.43×10^{-1}	1.78×10^{-1}
FNO	5.30×10^{-2}	8.20×10^{-2}	1.11×10^{-1}
DiffusionPDE	7.31×10^{-1}	2.44×10^{-1}	1.02×10^0
sCM-PINN	2.80×10^{-2}	8.91×10^{-2}	1.36×10^{-1}

Table 8: Inverse problem results (reconstructing initial state w_0) on the Navier-Stokes system. We report the relative H^1 error.

Method	NFE	Inverse H^1
DiffusionPDE	31	1.33×10^0
	63	1.17×10^0
	127	9.17×10^{-1}
sCM-PINN (Ours)	17	8.17×10^{-1}
	33	7.24×10^{-1}
	65	6.68×10^{-1}

Table 9: Distributional metrics and PDE residuals (unconditional sampling) on the Navier-Stokes system across varying computational budgets.

Method	NFE	$\ \mathcal{R}\ _2 \cdot h^2$	WD	JS
DiffusionPDE	63	1.92×10^{-2}	4.45×10^{-3}	3.45×10^{-4}
sCM	2	2.12×10^{-2}	2.40×10^{-2}	2.97×10^{-3}
	17	1.93×10^{-2}	2.08×10^{-3}	1.23×10^{-4}
sCM-PINN (Ours)	2	7.52×10^{-3}	1.24×10^{-1}	1.24×10^{-1}
	17	1.87×10^{-2}	2.81×10^{-3}	1.42×10^{-4}

E.2 Distributional Fidelity and Concurrent Baselines

To demonstrate that sCM-PINN ensures distributional fidelity alongside physical accuracy, we benchmark against concurrent physics-informed generative models (PBFM and PIDM). For this evaluation, we utilize the specific Darcy Flow dataset introduced by PBFM and PIDM, which is distinct from the primary Darcy Flow dataset analyzed in the main paper. To ensure a strictly fair comparison and to address reporting discrepancies regarding generative quality, we evaluated all models using PBFM’s exact residual implementation alongside our distributional metrics.

We evaluate distributional fidelity using the Wasserstein distance (WD) and Jensen-Shannon divergence (JS), while physical fidelity is measured via the normalized PDE residual $\|\mathcal{R}\|_2 \cdot h^2$. During training, sCM-PINN utilized the exact same network hyperparameters from our Helmholtz and Poisson setups, demonstrating its out-of-the-box robustness. Additionally, for the physics-informed fine-tuning on this specific dataset, we set the residual loss norm to $p = 1$ (as defined in Eq. 10).

As shown in Table 7, sCM-PINN drastically reduces PDE residuals while maintaining WD (1.93×10^{-2}) and JS (7.23×10^{-5}) orders of magnitude better than PBFM and PIDM at just 2 NFE (vs. 20-100 NFE).

Table 7: Distributional metrics and physical fidelity on the Darcy Flow dataset.

Model	NFE	$\ \mathcal{R}\ _2 \cdot h^2$	WD	JS
sCM	2	4.26×10^{-3}	1.85×10^{-2}	7.56×10^{-5}
sCM-PINN	2	2.39×10^{-3}	1.93×10^{-2}	7.23×10^{-5}
PBFM	20	2.11×10^{-4}	2.70×10^{-1}	1.72×10^{-2}
PIDM	100	2.98×10^{-4}	9.83×10^{-2}	1.33×10^{-2}

E.3 Extended Evaluation on the Navier-Stokes System

To adapt our generative formulation to the Navier-Stokes system, we map the initial vorticity state w_0 to the coefficient channel ($\mathbf{a} = w_0$) and the final state w_T to the solution channel ($\mathbf{u} = w_T$). By jointly modeling the distribution of (w_0, w_T) , the frozen-decoder strategy generalizes effectively, yielding a unified inference framework capable of solving both the forward problem (predicting w_T given w_0) and the inverse problem (reconstructing w_0 from w_T).

Tables 8 and 9¹ summarize our findings. For the inverse problem, sCM-PINN outperforms DiffusionPDE’s H^1 error in fewer steps. In the unconditional generative setting, sCM-PINN achieves the lowest overall PDE residual at just 2 NFE. Furthermore, when scaled to 17 NFE, our model surpasses the 63-step DiffusionPDE baseline across all evaluated metrics (PDE residual, WD, and JS), while achieving distributional fidelity (WD/JS) highly comparable to the base sCM.

¹Distributional metrics differ from our OpenReview rebuttal due to a corrected evaluation implementation.

E.4 Sensitivity and Training Cost

As shown in Table 10 for the Darcy Flow benchmark, transitioning to Stage 2 at varying Stage 1 checkpoints (epochs 4, 6, and 8) yields stable metrics within the same order of magnitude. This demonstrates that our pipeline is robust to transition timing without requiring precise scheduling. Total training takes approximately 17 hours on a single NVIDIA RTX 4090 GPU. Finally, for all evaluations, we standardized our inference budget to 65 NFE for forward problems and 2 NFE for unconditional sampling.

Table 10: Sensitivity to Stage 1 to Stage 2 transition epoch (Darcy Flow).

Stage 1 epoch	Relative H^1 (forward)	$\ \mathcal{R}\ _2 \cdot h^2$ (unconditional)
4	1.08×10^{-1}	2.35×10^{-2}
6	1.09×10^{-1}	2.43×10^{-2}
8	1.05×10^{-1}	2.00×10^{-2}