

Multi-view Consistency and Frequency-aware Modeling for Scattering Scene Reconstruction

Renrong Hu^{1,2}, Qianye He¹, Dongyu Du^{1,2}, Yihui Fan¹, Zhiheng Li^{2,1}, Xin Jin^{1,*}

¹ Shenzhen International Graduate School, Tsinghua University, Shenzhen, China

² Tsinghua Innovation Center in Zhuhai, Zhuhai, China

Abstract—Reconstructing complex scenes in participating media, such as foggy or underwater environments, remains challenging due to severe light scattering and absorption. Recent hybrid representations that combine explicit surface modeling and implicit volumetric rendering have shown promising results, yet they still suffer from two critical issues: multi-view inconsistency caused by view-dependent scattering, leading to unstable color and geometry across viewpoints; and frequency-domain interference between explicit surface and implicit volumetric representations, which results in geometric distortion, over-smoothed details, and optimization instability. In this paper, we address these problems from a signal-processing perspective and propose a unified method for multi-view consistency modeling and frequency-aware representation decoupling in scattering scenes. First, we introduce a multi-view scattering consistency model that jointly enforces color and geometric coherence across viewpoints. Specifically, we extend haze-lines from single-image enhancement to a multi-view color subspace alignment mechanism, and incorporate a Microflake-based directional scattering prior to constrain view-invariant volumetric scattering behavior. We further introduce a frequency-aware representation decoupling strategy that separates low-frequency volumetric scattering from high-frequency surface geometry, mitigating frequency aliasing during joint optimization. Experiments on synthetic and real-world datasets show stable reconstruction-quality improvements over representative baselines. More importantly, the proposed method reduces Δ PSNR, color variance, and depth variance by 14.8%, 27.3%, and 25.4%, respectively, indicating improved cross-view stability under scattering.

Index Terms—Scattering scene reconstruction, participating media, multi-view consistency, hybrid scene representation, 3D Gaussian Splatting.

I. INTRODUCTION

Reconstructing and rendering scenes in participating media, such as fog, haze, and turbid water, is a fundamental problem in computational imaging and vision. In these environments, light undergoes complex absorption and scattering processes before reaching the camera, resulting in degraded image quality, color distortion, and unreliable geometric cues. These effects severely challenge traditional multi-view reconstruction and novel view synthesis methods, which typically assume clear atmospheric conditions.

Recent advances in neural scene representations have improved reconstruction quality under adverse conditions. Implicit volumetric representations, such as Neural Radiance Fields (NeRF) [1], are effective at modeling global light transport and volumetric scattering, while explicit surface-based representations, such as 3D Gaussian Splatting (3DGS) [2],

excel at capturing fine geometric details and high-frequency textures. Motivated by their complementary strengths, hybrid approaches that combine explicit surface modeling with implicit volumetric rendering have emerged as a promising direction for scene reconstruction and novel view synthesis [22]–[24]. However, despite their empirical success, existing hybrid models still exhibit notable limitations when applied to multi-view scattering scenes.

Despite these advances, two fundamental challenges remain unresolved. First, multi-view inconsistency persists due to view-dependent scattering effects, including varying optical path lengths and anisotropic scattering [4]. As a result, the same 3D point may exhibit inconsistent color, brightness, and even depth estimates across different viewpoints [5]–[7], leading to color drift, geometric instability, and view-dependent artifacts that degrade both reconstruction accuracy and rendering fidelity [8], [9]. Most existing scattering-aware methods address degradation in a per-view manner and lack explicit mechanisms to enforce cross-view coherence of color, geometry, and scattering behavior.

Second, hybrid scene representations suffer from frequency-domain conflicts during joint optimization [10]. Explicit surface models inherently emphasize high-frequency structures such as sharp edges and fine textures, whereas implicit volumetric models predominantly capture low-frequency components, including smooth attenuation, global illumination, and haze-like effects [11]–[13]. Without explicit frequency allocation, these representations compete for overlapping signal components, resulting in frequency aliasing [14]. Consequently, low-frequency scattering effects may be incorrectly absorbed into surface geometry [15], [16], while high-frequency details are over-smoothed by volumetric modeling, ultimately degrading reconstruction quality and optimization stability [20], [21].

In this work, we propose a unified hybrid method that explicitly enforces cross-view scattering consistency and reduces frequency-domain conflicts between explicit surface and implicit volumetric representations. Unlike hybrid methods that mainly improve the representation backbone, our method focuses on two failure modes that are particularly common in participating media: view-dependent color and geometry drift, and ambiguous frequency allocation between volumetric scattering and surface details.

Our contributions are summarized as follows:

- We formulate multi-view scattering reconstruction as a consistency-oriented optimization problem, where recon-

*Corresponding author

struction quality is evaluated together with cross-view color and depth stability.

- We introduce haze-lines color consistency and Microflake scattering consistency into a unified multi-view optimization objective to constrain color, scattering, and geometry across viewpoints.
- We propose frequency-aware representation decoupling to assign low-frequency volumetric effects and high-frequency surface details to different branches of the hybrid representation.

The remainder of this paper is organized as follows. Section II presents the proposed method, including multi-view scattering consistency modeling and frequency-aware representation decoupling. Section III describes the experimental setup and provides quantitative and qualitative evaluations. Finally, Section IV concludes the paper and discusses future work.

II. METHOD

A. Overall Framework

We address the problem of multi-view scene reconstruction in participating media. Given a set of calibrated views

$$\mathcal{V} = \{(I_v, \mathbf{P}_v)\}_{v=1}^N, \quad (1)$$

where $I_v \in \mathbb{R}^{H \times W \times 3}$ denotes the observed image and $\mathbf{P}_v$ denotes the corresponding camera parameters, our goal is to reconstruct a scene representation that supports geometrically stable reconstruction and color-consistent novel-view rendering under scattering.

As illustrated in Fig. 1, our framework consists of a hybrid representation backbone and a set of explicit multi-view consistency constraints. The backbone combines a NeRF-based volumetric scattering representation and a 3D Gaussian Splatting (3DGS) surface representation to jointly capture low-frequency volumetric effects and high-frequency surface details. Built upon this representation, three complementary components are introduced to enforce cross-view color coherence, scattering stability, and frequency-domain separation. The method does not rely on a single isolated module as its contribution; instead, it combines physically motivated color and scattering constraints with frequency allocation so that the hybrid representation is optimized toward stable multi-view behavior in participating media.

In participating media, the observed color of a 3D point $\mathbf{x}$ from view v can be expressed as

$$\mathbf{I}_v(\mathbf{x}) = \mathbf{J}(\mathbf{x}) t_v(\mathbf{x}) + \mathbf{A}_v(1 - t_v(\mathbf{x})), \quad (2)$$

where $\mathbf{J}(\mathbf{x})$ denotes the underlying surface radiance, $\mathbf{A}_v$ is the ambient scattering (airlight) color, and

$$t_v(\mathbf{x}) = \exp\left(-\int_0^{d_v(\mathbf{x})} \sigma(s) ds\right) \quad (3)$$

is the view-dependent transmittance determined by the extinction coefficient $\sigma(\cdot)$. This formulation highlights the inherent view dependence of both color and geometry under scattering, motivating explicit multi-view constraints.

B. Hybrid Representation Backbone

1) *3D Gaussian Splatting for Surface Modeling*: We employ 3D Gaussian Splatting (3DGS) to explicitly model scene surfaces. Each Gaussian primitive encodes spatial position, covariance, opacity, and appearance attributes, enabling efficient rasterization and accurate reconstruction of high-frequency surface details such as object boundaries and textures. While effective for surface geometry, explicit surface representations are sensitive to view-dependent scattering artifacts when applied to participating media.

2) *NeRF-based Volumetric Scattering Modeling*: To model volumetric scattering effects, we adopt a NeRF-based volumetric representation. Given spatial coordinates and viewing directions, the volumetric branch predicts medium density and scattering-aware color, which are integrated along camera rays via volume rendering. This branch primarily captures low-frequency volumetric effects such as absorption, in-scattering, and out-scattering, but may suffer from view-dependent ambiguity without additional constraints.

C. Multi-view Consistency Constraints

1) *Haze-Lines Color Consistency*: For a fixed 3D point $\mathbf{x}$ observed from multiple views, haze-lines theory states that the observed colors approximately lie on a line in RGB space connecting the clean color and the ambient scattering color. We model this relationship as

$$\mathbf{I}_v(\mathbf{x}) \approx \mathbf{A}_v + \alpha_v(\mathbf{x})(\mathbf{J}(\mathbf{x}) - \mathbf{A}_v), \quad (4)$$

where $\alpha_v(\mathbf{x}) \in [0, 1]$ is a view-dependent mixing coefficient related to transmittance.

We enforce a shared latent clean color $\mathbf{J}(\mathbf{x})$ across all views and define the haze-lines consistency loss as

$$\mathcal{L}_{\text{HL}} = \sum_{\mathbf{x}} \sum_v \|\mathbf{I}_v(\mathbf{x}) - \mathbf{A}_v - \alpha_v(\mathbf{x})(\mathbf{J}(\mathbf{x}) - \mathbf{A}_v)\|_2^2. \quad (5)$$

This constraint restricts cross-view color variations to a physically interpretable one-dimensional subspace, effectively suppressing view-dependent color drift.

2) *Microflake Scattering Consistency*: While haze-lines enforce color coherence, geometric instability often arises from view-dependent volumetric scattering. To address this issue, we adopt a Microflake-based scattering model, which represents the medium as a collection of oriented microstructures inducing anisotropic scattering. Let $p(\omega_i \rightarrow \omega_o \mid \boldsymbol{\theta}(\mathbf{x}))$ denote the Microflake phase function parameterized by $\boldsymbol{\theta}(\mathbf{x})$, encoding local medium properties.

The scattered radiance at $\mathbf{x}$ can be written as

$$L_o(\mathbf{x}, \omega_o) = \int p(\omega_i \rightarrow \omega_o \mid \boldsymbol{\theta}(\mathbf{x})) L_i(\mathbf{x}, \omega_i) d\omega_i. \quad (6)$$

Since the medium properties at a fixed spatial location should be view-invariant, we enforce directional scattering consistency across all views observing $\mathbf{x}$:

$$\begin{aligned} \mathcal{L}_{\text{MF}} &= \sum_{\mathbf{x}} \sum_{v \in \mathcal{V}_{\mathbf{x}}} \|L_v(\mathbf{x}) - \bar{L}(\mathbf{x})\|_2^2, \\ \bar{L}(\mathbf{x}) &= \frac{1}{|\mathcal{V}_{\mathbf{x}}|} \sum_{v \in \mathcal{V}_{\mathbf{x}}} L_v(\mathbf{x}), \end{aligned} \quad (7)$$

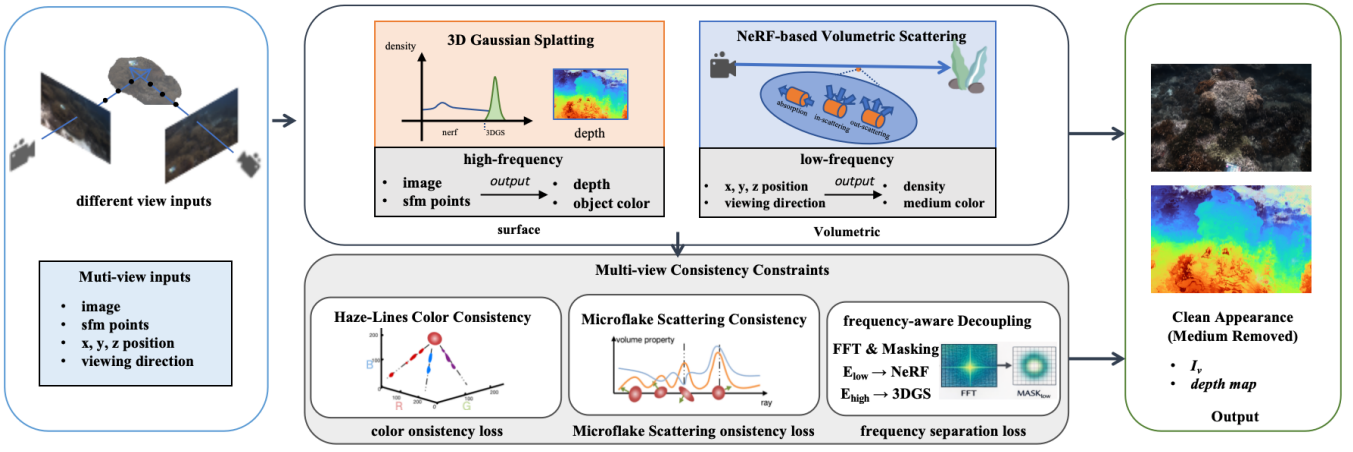


Fig. 1. Overview of the proposed hybrid representation framework for multi-view scene reconstruction in participating media. The framework employs a hybrid backbone that combines NeRF-based volumetric scattering modeling for low-frequency effects and 3D Gaussian Splatting for high-frequency surface geometry. Built upon this representation, three coordinated components are introduced: (1) haze-lines color consistency, which enforces view-invariant color subspace alignment; (2) Microflake-based scattering consistency, which stabilizes volumetric scattering behavior and geometry across viewpoints; and (3) frequency-aware representation decoupling, which explicitly separates low- and high-frequency components and assigns them to volumetric and surface modeling, respectively. Together, these components improve geometric stability and color consistency under complex scattering conditions.

where $\mathcal{V}_x$ denotes the set of views in which x is visible. This constraint stabilizes volumetric scattering behavior and indirectly improves geometric consistency.

3) *Frequency-aware Representation Decoupling*: Explicit surface and volumetric representations exhibit distinct frequency characteristics. Surface-based models concentrate energy in high-frequency bands, whereas volumetric models dominate low-frequency components. Without explicit separation, joint optimization leads to frequency aliasing between the two representations.

Given an intermediate feature representation $\mathbf{E}$, we perform frequency decomposition using FFT-based masking:

$$\mathbf{E}^{\text{low}} = \mathcal{F}^{-1}(\mathbf{M}_{\text{low}} \odot \mathcal{F}(\mathbf{E})), \quad (8)$$

$$\mathbf{E}^{\text{high}} = \mathcal{F}^{-1}(\mathbf{M}_{\text{high}} \odot \mathcal{F}(\mathbf{E})), \quad (9)$$

where $\mathbf{M}_{\text{low}}$ and $\mathbf{M}_{\text{high}}$ are complementary frequency masks.

We explicitly allocate low-frequency components to the volumetric branch and high-frequency components to the surface branch. To further suppress frequency leakage, we introduce a frequency regularization term:

$$\mathcal{L}_{\text{FR}} = \left\| \mathcal{F}(\hat{I}_{\text{surf}}) \odot \mathbf{M}_{\text{low}} \right\|_2^2 + \left\| \mathcal{F}(\hat{I}_{\text{vol}}) \odot \mathbf{M}_{\text{high}} \right\|_2^2. \quad (10)$$

D. Overall Objective

The final training objective is defined as

$$\mathcal{L} = \mathcal{L}_{\text{photo}} + \lambda_{\text{HL}} \mathcal{L}_{\text{HL}} + \lambda_{\text{MF}} \mathcal{L}_{\text{MF}} + \lambda_{\text{FR}} \mathcal{L}_{\text{FR}}, \quad (11)$$

where $\mathcal{L}_{\text{photo}}$ denotes the standard photometric reconstruction loss, and λ_{HL} , λ_{MF} , and λ_{FR} are weighting coefficients.

III. EXPERIMENTS

A. Experimental Setup

All experiments are conducted on a workstation equipped with six NVIDIA RTX A40 GPUs (48GB memory per GPU).

Our implementation is based on PyTorch and optimized using distributed data parallel training. Unless otherwise specified, all models are trained using the same optimization settings to ensure fair comparison. We use the Adam optimizer with an initial learning rate of 5×10^{-4} , which is exponentially decayed during training. The model is trained until convergence with early stopping based on validation loss. Loss weights are empirically set to $\lambda_{\text{HL}} = 1.0$, $\lambda_{\text{MF}} = 0.5$, and $\lambda_{\text{FR}} = 0.1$ unless stated otherwise. All experiments adopt mixed-precision training to improve computational efficiency.

B. Dataset

We evaluate our method on a multi-view scattering dataset adopted from [20], which follows established data acquisition protocols for scattering imaging and underwater vision. The dataset consists of multiple static scenes captured under complex participating media, including both foggy atmospheric conditions and underwater environments, where light undergoes severe absorption and multiple scattering effects.

Specifically, the dataset includes representative aerial urban scenes under dense fog as well as underwater industrial scenes with varying water turbidity. As illustrated in Fig. 2, the foggy scenes feature large-scale outdoor environments with roads, buildings, and vegetation observed from aerial viewpoints, while the underwater scenes contain man-made structures and equipment captured in real underwater conditions. These scenarios pose challenges for both geometric reconstruction and appearance consistency due to strong view-dependent scattering.

For each scene, 20 to 50 viewpoints are captured using calibrated cameras, providing accurate camera poses and multi-view RGB observations. The dataset covers a wide range of scattering densities, visibility levels, and object materials, enabling comprehensive evaluation across diverse scattering

TABLE I
QUANTITATIVE COMPARISON ON THE SELF-COLLECTED SCATTERING DATASET.

Method	PSNR (dB)↑	SSIM↑	LPIPS↓
NeRF [1]	26.412	0.701	0.312
3DGS [2]	28.315	0.856	0.269
ScatterSplatting [22]	31.164	0.892	0.218
Ours	31.685	0.924	0.197

TABLE II
MULTI-VIEW CONSISTENCY EVALUATION ON THE SCATTERING DATASET.
LOWER VALUES INDICATE BETTER CROSS-VIEW STABILITY.

Method	Δ PSNR↓	Color Var.↓	Depth Var.↓
NeRF [1]	1.92	1.89e-2	3.94e-2
3DGS [2]	1.41	1.72e-2	3.15e-2
ScatterSplatting [22]	1.08	1.21e-2	2.64e-2
Ours	0.92	0.88e-2	1.97e-2

conditions. Following [20], both RGB images and corresponding camera parameters are provided for all scenes.

C. Baselines

We compare the proposed method with representative approaches covering implicit volumetric modeling, explicit surface reconstruction, and hybrid representations. Specifically, we consider standard NeRF [1] trained under clear media assumptions, 3D Gaussian Splatting (3DGS) [2] for explicit surface reconstruction, and a hybrid NeRF+3DGS model ScatterSplatting [22] trained without explicit multi-view consistency or frequency-aware constraints. All baselines are re-trained on our dataset using their official implementations and optimized for best performance. Recent hybrid representations, such as NeRF-GS [23] and HyRF [24], also combine implicit and explicit scene modeling. Our comparison focuses on representative baselines available under the same scattering-scene setting, while the proposed constraints are orthogonal to the choice of hybrid backbone and target the scattering-specific consistency problem that these general hybrid methods do not explicitly address.

D. Results

Table I reports the average reconstruction quality on the scattering dataset. Our method achieves the best overall performance, obtaining the highest PSNR and SSIM and the lowest LPIPS. Compared with ScatterSplatting [22], our method improves PSNR from 31.164 dB to 31.685 dB, corresponding to a gain of 0.521 dB, while also improving SSIM from 0.892 to 0.924 and reducing LPIPS from 0.218 to 0.197. These results indicate that the proposed multi-view consistency modeling and frequency-aware representation decoupling improve both pixel-level fidelity and perceptual quality beyond the hybrid NeRF-3DGS backbone.

However, average image-quality metrics alone are insufficient to characterize reconstruction stability in participating media, where view-dependent scattering may cause color

drift, depth ambiguity, and viewpoint-specific degradation. The relatively moderate PSNR gain further motivates the need to evaluate cross-view stability, where scattering-specific improvements are more directly reflected. Therefore, we further report consistency-related metrics in Table II. Compared with ScatterSplatting, our method reduces Δ PSNR from 1.08 to 0.92, Color Var. from $1.21e-2$ to $0.88e-2$, and Depth Var. from $2.64e-2$ to $1.97e-2$, corresponding to relative reductions of 14.8%, 27.3%, and 25.4%, respectively. These improvements show that our method not only improves per-view reconstruction quality, but also produces more stable appearance and geometry across viewpoints.

Among the evaluated baselines, NeRF exhibits the largest cross-view variations, indicating that implicit volumetric modeling alone suffers from severe view-dependent instability in scattering scenes. 3DGS improves geometric sharpness through explicit surface modeling, but cannot fully account for volumetric scattering effects. ScatterSplatting further improves stability by combining surface and volumetric representations, yet still lacks explicit cross-view scattering constraints. In contrast, our method achieves the lowest color and depth variations, demonstrating the effectiveness of haze-lines based color subspace alignment and the Microflake-based directional scattering prior.

The quantitative results are consistent with the qualitative comparisons in Fig. 2. In foggy aerial scenes, competing methods retain noticeable color shifts in distant regions and blurred boundaries around high-contrast structures. In underwater scenes, residual scattering and low-visibility regions lead to unstable appearance and geometry. Our method reduces these artifacts and produces more coherent colors, sharper structures, and more stable geometry across viewpoints. Overall, the results demonstrate that explicitly modeling multi-view scattering consistency and frequency-aware representation decoupling improves both reconstruction fidelity and cross-view stability in complex scattering scenes.

E. Efficiency Analysis

We also evaluate the computational overhead introduced by the proposed constraints. The haze-lines, Microflake, and frequency-aware terms are implemented as training-time regularization and do not introduce additional inference branches. As a result, the training time increases only from 8.6 minutes to 8.9 minutes on the same experimental setting, corresponding to an overhead of approximately 3%. The inference speed remains 42.5 FPS, matching the hybrid backbone and indicating that the consistency gains are obtained without extra test-time cost.

F. Ablation Study

To analyze the contribution of each component, we conduct an ablation study by selectively removing individual modules. Table III summarizes the quantitative results.

Removing the haze-lines color consistency constraint leads to noticeable performance degradation, primarily reflected by increased LPIPS and reduced PSNR. More importantly, the

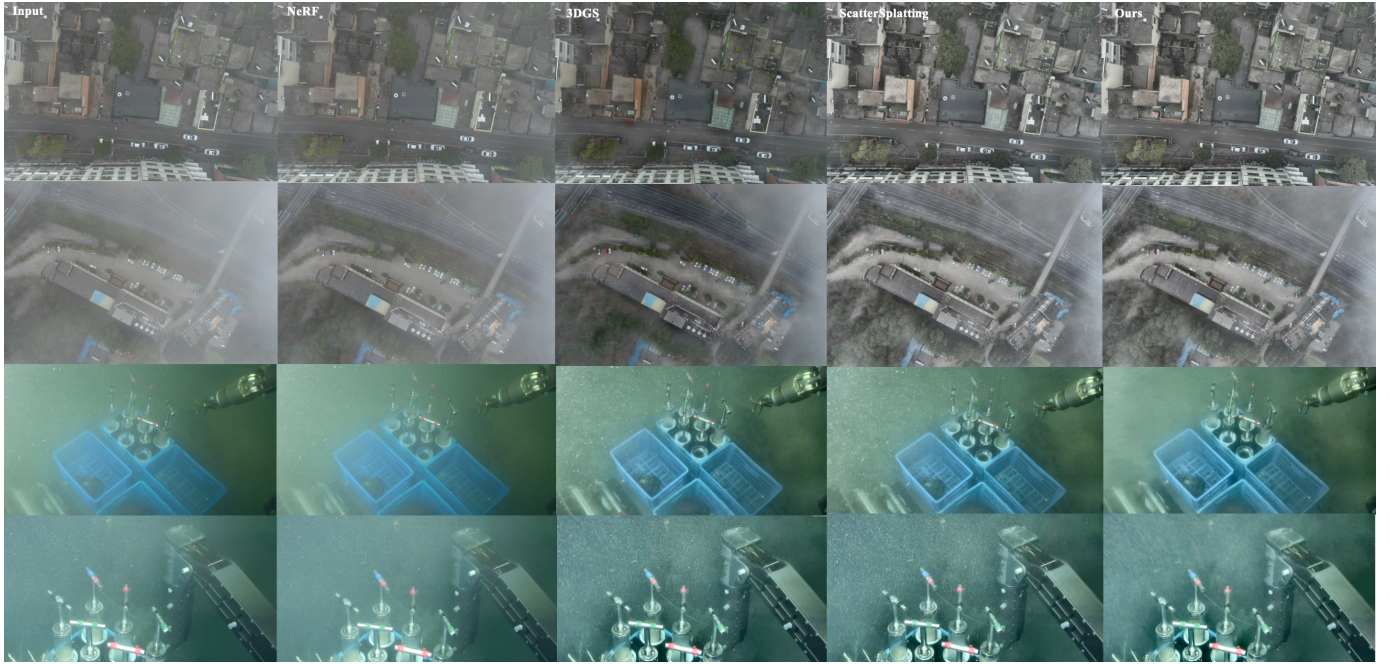


Fig. 2. Qualitative comparison on representative foggy and underwater scattering scenes. Columns show the input degraded image, NeRF optimization [1], 3D Gaussian Splatting (3DGS) optimization [2], ScatterSplatting optimization [22], and our method. The comparison highlights view-dependent color drift, residual scattering, and structural instability in challenging regions such as distant backgrounds, object boundaries, and low-visibility underwater areas.

TABLE III
ABLATION STUDY OF DIFFERENT COMPONENTS.

Configuration	PSNR $\uparrow$	SSIM $\uparrow$	LPIPS $\downarrow$
Full model	31.685	0.924	0.197
w/o HL	30.982	0.881	0.226
w/o MF	30.147	0.873	0.241
w/o FR	31.184	0.885	0.219

absence of haze-lines consistency increases cross-view color variation, indicating view-dependent color drift.

When the Microflake scattering consistency is removed, geometric stability degrades more noticeably. Although the average image quality metrics decrease moderately, depth predictions become highly inconsistent across views, demonstrating that enforcing direction-aware, physically consistent scattering behavior is crucial for stabilizing geometry under participating media.

Disabling frequency-aware representation decoupling results in consistent but milder performance drops. Compared with the full model, removing this component decreases PSNR from 31.685 dB to 31.184 dB, an approximately 1.6% reduction. This drop is consistent with frequency aliasing artifacts and reduced optimization stability, showing that frequency-aware decoupling complements multi-view consistency constraints by reducing ambiguity between volumetric scattering and surface details.

IV. CONCLUSION

In this paper, we proposed a unified hybrid method for multi-view scene reconstruction in participating media that jointly enforces multi-view scattering consistency and frequency-aware representation decoupling. By extending haze-lines to a cross-view color subspace alignment mechanism and incorporating a Microflake-based directional scattering prior, the method stabilizes color, geometry, and volumetric scattering behavior across viewpoints. In addition, a frequency-aware allocation strategy is introduced to disentangle low-frequency volumetric effects from high-frequency surface geometry, reducing frequency interference during joint optimization.

Experiments show consistent reconstruction-quality improvements and clearer gains in cross-view color and geometric stability, while maintaining the inference speed of the hybrid backbone. The current evaluation focuses on static scenes captured with 20 to 50 calibrated views; future work will study fewer-view settings, dynamic scenes, and adaptive frequency separation strategies.

ACKNOWLEDGMENT

This work was supported in part by NSF of Guangdong Province under Grant 2023A1515012716 and Shenzhen Science and Technology Program under the Grant KCXFZ20240903094301003.

REFERENCES

- [1] Mildenhall B, Srinivasan P P, Tancik M, et al. NeRF: Representing scenes as neural radiance fields for view synthesis[J]. Communications of the ACM, 2021, 65(1): 99–106.

- [2] Kerbl B, Kopanas G, Leimkühler T, et al. 3D Gaussian Splatting for Real-Time Radiance Field Rendering[J]. *ACM Transactions on Graphics*, 2023, 42(4): 139:1–139:14.
- [3] Barron J T, Mildenhall B, Verbin D, et al. Mip-NeRF 360: Unbounded Anti-Aliased Neural Radiance Fields[C]// *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2022: 5470–5479.
- [4] Mildenhall B, Hedman P, Martin-Brualla R, et al. NeRF in the Dark: High Dynamic Range View Synthesis from Noisy Raw Images[C]// *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2022: 16190–16199.
- [5] Levy D, Peleg A, Pearl N, et al. SeaThru-NeRF: Neural Radiance Fields in Scattering Media[C]// *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2023: 56–65.
- [6] Tang Y, Zhu C, Wan R, et al. Neural Underwater Scene Representation[C]// *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2024: 11780–11789.
- [7] Li T, Li L U, Wang W, et al. Dehazing-NeRF: Neural Radiance Fields from Hazy Images[J]. *arXiv preprint arXiv:2304.11448*, 2023.
- [8] Yu Z, Chen A, Huang B, et al. Mip-Splatting: Alias-Free 3D Gaussian Splatting[C]// *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2024: 19447–19456.
- [9] Yu Z, Sattler T, Geiger A. Gaussian Opacity Fields: Efficient and Compact Surface Reconstruction in Unbounded Scenes[J]. *arXiv preprint arXiv:2404.10772*, 2024.
- [10] Ye Z, Li W, Liu S, et al. AbsGS: Recovering Fine Details in 3D Gaussian Splatting[C]// *Proceedings of ACM Multimedia*, 2024.
- [11] Akkaynak D, Treibitz T. A Revised Underwater Image Formation Model[C]// *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2018: 6723–6732.
- [12] Sakaridis C, Dai D, Van Gool L. Guided Curriculum Model Adaptation and Uncertainty-Aware Evaluation for Semantic Nighttime Image Segmentation[C]// *Proceedings of the IEEE/CVF International Conference on Computer Vision*, 2019: 7374–7383.
- [13] Mildenhall B, Srinivasan P P, Ortiz-Cayon R, et al. Local Light Field Fusion: Practical View Synthesis with Prescriptive Sampling Guidelines[J]. *ACM Transactions on Graphics*, 2019, 38(4): 1–14.
- [14] Berman D, Treibitz T, Avidan S. Single Image Dehazing Using HazeLines[J]. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2018, 42(3): 720–734.
- [15] Zhang K, Bi S, Sunkavalli K, et al. Neural Microfacet Fields for Inverse Rendering[C]// *Proceedings of the IEEE/CVF International Conference on Computer Vision*, 2023.
- [16] Marschner S R, Hanrahan P. Practical Modeling and Rendering of Structured Hair[J]. *ACM Transactions on Graphics*, 2003, 22(3): 780–791.
- [17] Liu Y, Wu Q, Xu D, et al. FreeNeRF: Improving Few-Shot Neural Rendering with Free Frequency Regularization[C]// *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2023.
- [18] Xu H, Zhang J, Chen Z, et al. FreGS: Frequency Regularized Gaussian Splatting for Scene Reconstruction[C]// *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2024.
- [19] Pharr M, Jakob W, Humphreys G. *Physically Based Rendering: From Theory to Implementation*[M]. Morgan Kaufmann, 2016.
- [20] Fan Y, Li Y, Zhang K, et al. Converting a Conventional Camera to a Super-Camera: Directional Atmospheric Scattering Modeling for Passive Imaging in Intense Real-World Scattering Scenarios[J]. *Advanced Photonics Nexus*, 2025, 4(5): 056004.
- [21] Fan Y, Li Y, Zhang K, et al. Depth-Rectified Statistical Scattering Modeling for Deep-Sea Video Descattering[J]. *Infrared and Laser Engineering*, 2022, 51(9): 20210919.
- [22] R. Hu, Q. He, D. Du and X. Jin, "ScatterSplatting: Enhanced View Synthesis in Scattering Scenarios via Joint NeRF and Gaussian Splatting," 2025 IEEE International Symposium on Circuits and Systems (ISCAS), London, United Kingdom, 2025, pp. 1-5, doi: 10.1109/ISCAS56072.2025.11044088.
- [23] Fang S, Shen I-C, Igarashi T, et al. NeRF Is a Valuable Assistant for 3D Gaussian Splatting[C]// *Proceedings of the IEEE/CVF International Conference on Computer Vision*, 2025: 26230–26240.
- [24] Wang Z, Xu D. HyRF: Hybrid Radiance Fields for Memory-efficient and High-quality Novel View Synthesis[J]. *arXiv preprint arXiv:2509.17083*, 2025.