

RetinexGS: A 3D Retinex Framework for Gaussian Splatting in Low-Light

Xingzhe Luo¹, Yingbo Wu^{1,2,*}, Wenxin Li¹, Haoran Wang¹

¹National Elite Institute of Engineering, Chongqing University, Chongqing, China

²School of Big Data & Software Engineering, Chongqing University, Chongqing, China

wyb@cqu.edu.cn

Abstract—Reconstructing and rendering high-fidelity, normally-lit novel views from multi-view low-light images via 3D Gaussian Splatting remains a key challenge. Low signal-to-noise ratio and color distortion of input images severely degrade geometric reconstruction accuracy. Conventional 2D enhancement preprocessing breaks multi-view photometric consistency, causing geometric artifacts. Existing advanced 3D decomposition methods, which rely solely on weak unsupervised priors, cannot stably realize visually consistent reflectance-illumination disentanglement. To address these issues, we propose RetinexGS, a novel hybrid-supervised framework. It uses 2D Retinex decomposition as strong priors to guide 3DGS in disentangling scenes into view-consistent 3D reflectance and structure-aware illumination fields, with a jointly trained enhancement network enabling automatic brightness restoration. Notably, it reconstructs normally-lit, 3D-consistent, and photorealistic views from challenging low-light inputs. Experiments on a challenging benchmark show RetinexGS outperforms or matches state-of-the-art methods in both quantitative metrics and visual quality.

Index Terms—Novel view synthesis, 3D Gaussian Splatting, Low-light image enhancement, Retinex Decomposition

I. INTRODUCTION

Recent 3D scene representation methods like Neural Radiance Fields (Nerf) [1] and 3D Gaussian Splatting (3DGS) [2] have achieved groundbreaking progress in novel view synthesis, with 3DGS being particularly notable for enabling real-time, high-fidelity rendering. Subsequently, a series of methods [3]–[10] have been proposed to further improve its rendering quality. Yet these achievements rely on an ideal premise: input multi-view images captured under good, uniform lighting. In real-world scenarios, low-light environments are unavoidable—resulting in images with low brightness, heavy noise, and color distortion. Feeding these images into standard 3DGS leads to inaccurate geometric reconstruction, and rendered novel views inherit input defects, failing to produce normally-lit clear views.

Two types of methods address the challenges of 3DGS reconstruction and rendering in low-light scenarios: one involves preprocessing images with 2D low-light enhancement algorithms [11]–[14] before feeding them into the 3DGS pipeline—however, this 2D enhancement approach, which processes each image independently, causes multi-view photometric inconsistency, disrupts 3DGS’s learning of scene geometry and appearance. The other focuses on end-to-end scene



Fig. 1. Visual illustration of low-light enhancement. Given low-light input images (left), our method produces normally-lit outputs (right) with improved brightness, color fidelity, and structural details.

modeling in 3D space. Among these end-to-end methods, LLNeRF [15] and LLGS [16] adopt an intrinsic decomposition approach but rely on unsupervised priors to drive this ill-posed decomposition process, resulting in insufficient stability. In contrast, Aleth-nerf [17] and Luminance-GS [18] employ heuristic appearance modulation, which fails to achieve visually consistent intrinsic decoupling, limiting their potential for consistent appearance editing tasks.

To address these limitations, we propose RetinexGS, a hybrid-supervised 3DGS framework (see Figure 1). Rooted in Retinex theory [19]–[24], it uses mature 2D Retinex models to guide 3D learning—addressing both 3D-consistent Retinex application and unsupervised decomposition instability. Specifically, we first generate reflectance/illumination maps from low-light images via an advanced 2D Retinex model; despite their potential multi-view inconsistency, these maps provide essential initial guidance for 3D decomposition, allowing the model to converge towards a visually consistent state. We then integrate these maps as supervision into a new 3DGS loss function, guiding the model to decouple the 3D scene into view-consistent albedo and a structure-aware 3D illumination field. Additionally, a lightweight enhancement network takes decoupled 3D illumination as input, trained end-to-end via unsupervised perceptual loss to predict optimal enhancement parameters for brightness restoration. We emphasize that our decomposition is Retinex-inspired rather than strictly physical, and is designed to improve optimization stability and cross-view consistency for low-light novel view synthesis. Our main contributions are:

- We propose a novel hybrid-supervised paradigm for the intrinsic decomposition of 3D scenes.

*Corresponding author.

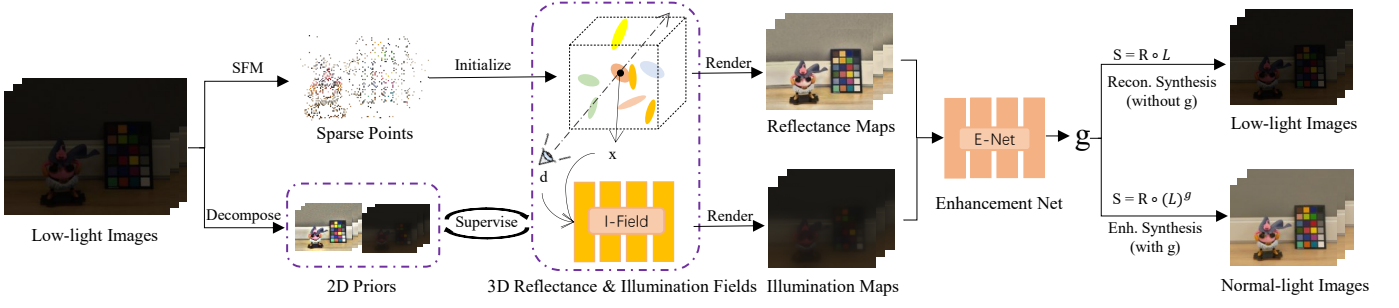


Fig. 2. Overview of RetinexGS. This framework leverages strong priors generated by 2D decomposition to supervise the learning of a hybrid 3D model for consistent low-light scene reconstruction. The optimization of the model is jointly driven by two synthesis pathways: the Reconstruction Synthesis pathway ensures the fidelity of decomposition under the supervision of 2D priors, while the Enhancement Synthesis pathway restores normal illumination and generates the final output via an Enhancement Network.

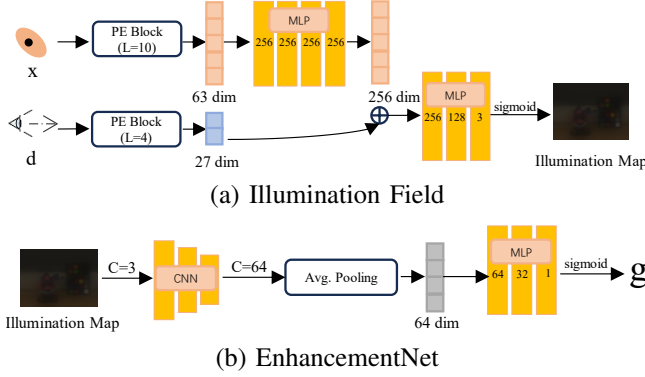


Fig. 3. Structure of the implicit illumination field and EnhancementNet. Numerical labels indicate each module’s output feature dimension.

- We implement an end-to-end trainable enhancement and reconstruction framework for low-light scenarios.
- We demonstrate that RetinexGS achieves highly competitive performance against state-of-the-art methods on the LOM dataset [17].

II. METHOD

A. Preliminary

3DGS models scenes explicitly using 3D Gaussian primitives with anisotropic properties, typically initialized from 3D point sets. Each primitive is defined by key parameters: a 3D center $\mu \in \mathbb{R}^3$, a covariance matrix $\Sigma \in \mathbb{R}^{3 \times 3}$, an opacity $\sigma \in \mathbb{R}^+$, and a view-dependent color c parameterized via spherical harmonic coefficients.

For rendering, 3D Gaussians are projected onto the image plane using camera pose (R, t) , yielding 2D Gaussians whose positions and covariance mappings are approximated. These 2D Gaussians are depth-sorted in camera space, and their contributions to each pixel are blended to compute the final color C :

$$C = \sum_{i=1}^n c_i \alpha_i \prod_{j=1}^{i-1} (1 - \alpha_j). \quad (1)$$

Here, $\alpha_i = \sigma_i G'_i(x')$, where $G'_i(x')$ is the value of the 2D projected Gaussian at pixel x' . Rendered results are compared

with ground truth for loss calculation to optimize Gaussian parameters.

B. 2D Decomposition for 3D Supervision

To address the ill-posed problem of unsupervised 3D scene decomposition, we introduce a hybrid supervision strategy that leverages a well-established 2D decomposition model to provide strong pixel-level priors for 3D learning. We adopt Lr3m [25] as the 2D decomposer. Lr3m is based on the classical Retinex theory, which models an image S as the product of its reflectance R and illumination L . Furthermore, to address common noise in low-light images, Lr3m employs a more robust model that includes a noise term N :

$$S = R \circ L + N. \quad (2)$$

This model effectively suppresses noise by incorporating low-rank priors, thereby generating clear reflectance maps.

In our framework, Lr3m serves as an offline preprocessing step. For each low-light input view S_i , we generate its corresponding 2D reflectance map R_i and illumination map L_i . Though these results have multi-view photometric inconsistency, they still offer a valuable initial estimation of the scene’s intrinsic properties. This set of 2D images $\{R_i, L_i\}$ acts as direct supervision signals to guide our 3D model in learning a visually plausible and globally consistent decomposition.

C. Hybrid 3D Scene Representation

Our core idea is to decouple the 3D scene into an explicit reflectance field and an implicit illumination field, as illustrated in our pipeline in Figure 2; this hybrid representation lays the foundation for subsequent intrinsic decomposition.

3D reflectance field. We leverage the explicit Gaussian primitives of 3DGS to model the view-independent intrinsic reflectance of the scene. The spherical harmonic coefficients of each Gaussian primitive are redefined to learn the diffuse albedo at that spatial point. During training, we regularize the high-order SH coefficients to ensure this field primarily represents the stable intrinsic properties of the scene, which is consistent with the definition of intrinsic diffuse albedo.

3D illumination field. To capture complex appearance effects that cannot be expressed by the reflectance field, we

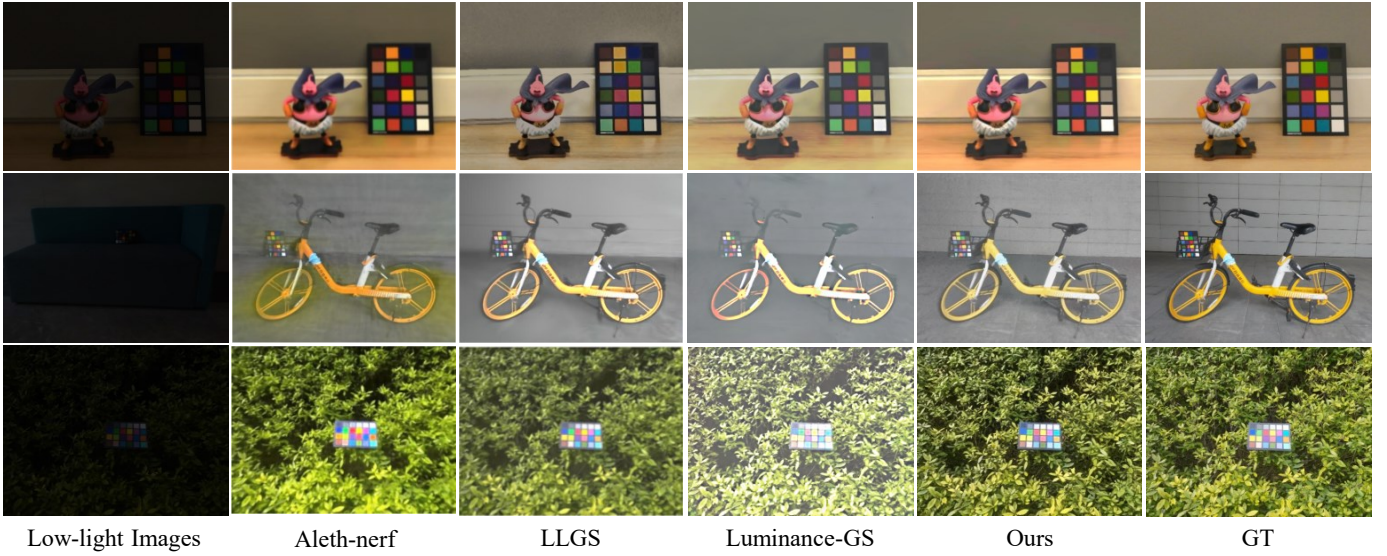


Fig. 4. Qualitative comparison with state-of-the-art methods on the LOM dataset.

TABLE I
THE QUANTITATIVE COMPARISON RESULTS ON THE LOM DATASET (PSNR $\uparrow$, SSIM $\uparrow$, LPIPS $\downarrow$). BEST RESULTS ARE BOLDDED.

Method	"buu"	"sofa"	"bike"	"chair"	"shrub"	mean
3DGS [2]	7.44/0.251/0.434	6.20/0.205/0.552	6.35/0.068/0.584	5.98/0.139/0.642	8.08/0.037/0.664	6.81/0.140/0.575
RetiNexNet [11]+3DGS	16.19/0.797/0.255	18.36/0.778/0.397	17.63/0.685/0.425	16.94/0.770/0.370	15.09/0.393/0.473	16.84/0.684/0.384
Zero-DCE [12]+3DGS	17.82/0.888/0.207	14.22/0.818/0.301	10.31/0.486/0.402	12.34/0.714/0.338	12.68/0.413/0.396	13.47/0.663/0.329
LLVE [13]+3DGS	19.85/0.865/0.261	17.82/0.842/0.382	13.14/0.587/0.496	15.09/0.764/0.445	15.76/0.377/0.459	16.33/0.687/0.409
SCI [14]+3DGS	7.79/0.694/0.411	10.16/0.718/0.354	13.28/0.561/0.374	20.37/0.823/0.346	17.79/0.564/0.407	13.88/0.663/0.378
3DGS+RetiNexNet [11]	17.84/0.756/0.387	20.08/0.775/0.485	17.53/0.628/0.543	18.24/0.718/0.515	16.55/0.465/0.514	18.05/0.681/0.489
3DGS+Zero-DCE [12]	17.17/0.826/0.313	14.37/0.787/0.423	10.37/0.471/0.515	12.51/0.686/0.469	12.85/0.393/0.443	13.45/0.632/0.433
3DGS+LLVE [13]	19.56/0.842/0.294	17.86/0.840/0.402	13.69/0.615/0.492	15.15/0.761/0.481	15.48/0.357/0.516	16.35/0.683/0.437
3DGS+SCI [14]	17.36/0.802/0.319	13.50/0.728/0.446	8.79/0.342/0.520	10.89/0.591/0.524	12.08/0.347/0.446	12.52/0.563/0.451
Aleth-nerf [17]	20.26/0.864/0.311	19.51/0.858/0.358	19.96/0.727/0.540	21.06/0.823/0.467	18.20/0.513/0.444	19.80/0.757/0.424
LLGS [16]	20.69/0.880/0.198	24.20/0.854/0.375	20.11/0.751/0.368	23.51/0.833/0.348	19.48/0.614/0.334	21.60/0.786/0.325
Luminance-GS [18]	21.08/0.914/0.175	25.34/0.923/0.221	17.48/0.748/0.403	17.53/0.848/0.326	11.65/0.615/0.276	18.68/0.810/0.280
RetinexGS	23.23/0.897/0.173	25.55/0.885/0.183	20.63/0.804/0.258	19.73/0.862/0.259	20.11/0.689/0.184	21.85/0.827/0.211

introduce a continuous Nerf-like implicit illumination field. While standard irradiance is view-independent, we adopt a generalized Retinex strategy by explicitly absorbing view-dependent effects into this component. This design prevents geometric overfitting by offloading high-frequency variations, effectively modeling complex radiance within the multiplicative Retinex structure. As shown in Figure 3, this network is an MLP that takes positionally encoded 3D coordinates x and observation direction d as inputs, and outputs the corresponding RGB illumination values. Positional encoding endows the network with the ability to fit high-frequency illumination details, while the continuity of the MLP ensures the smoothness of the illumination field. This compact model effectively models global illumination.

D. Model Optimization

We jointly optimize the decomposition and enhancement of the scene in an end-to-end manner through a multi-component composite loss function. The overall optimization objective L_{total} consists of two main parts: the decomposition loss L_{decomp} and the unsupervised enhancement loss L_{enhance} :

$$L_{\text{total}} = L_{\text{decomp}} + L_{\text{enhance}}. \quad (3)$$

Decomposition loss under hybrid supervision. To stably solve this ill-posed decomposition problem, our decomposition loss consists of a supervision loss L_{sup} and structural prior regularization terms. The supervision loss term aims to align the 3D decomposition results with 2D priors and original observations. Following the design of the original 3DGS, we adopt a hybrid supervision strategy that combines pixel-level

L1 loss and perceptual-level SSIM loss. The total supervision loss L_{sup} comprises the following three components:

$$L_{\text{sup}} = \lambda_r L_s(\hat{R}, R) + \lambda_l L_s(\hat{L}, L) + \lambda_{\text{rec}} L_s(\hat{R} \circ \hat{L}, S), \quad (4)$$

where L_s is a similarity function composed of weighted L1 and SSIM losses. The first and second terms of this formula constrain the rendered 3D reflectance map $\hat{R}$ and illumination map $\hat{L}$ to approximate the 2D priors R and L , respectively. The third term is a fidelity term, which ensures that the product of the decomposition results can accurately reconstruct the original low-light input image S . Here, λ_r , λ_l , and λ_{rec} are set to 0.5, 0.5, and 1, respectively.

In addition to these supervision terms, we also introduce two structural priors as regularizations. The reflectance smoothness prior L_{TV} imposes a total variation constraint on the rendered reflectance map $\hat{R}$ to encourage piecewise smoothness of material surfaces and suppress noise:

$$L_{\text{TV}}(\hat{R}) = \|\nabla \hat{R}\|_1. \quad (5)$$

Finally, the spherical harmonic coefficient sparsity regularization L_{SH} penalizes the energy of high-order spherical harmonic coefficients $c_{>0}$. This regularization serves as a soft constraint to strictly enforce view-independence, preventing the reflectance field from absorbing lighting effects. By suppressing high-order terms, it effectively constrains the reflectance to behave as a Lambertian surface and ensures that view-dependent radiance is offloaded to the illumination field, thereby enhancing the intrinsic consistency of the decomposition:

$$L_{\text{SH}} = \|c_{>0}\|_2^2. \quad (6)$$

The weights for both L_{TV} and L_{SH} are set to 0.01.

Unsupervised illumination enhancement. After decoupling the scene’s reflectance and illumination, we designed a lightweight EnhancementNet (structure is shown in Figure 3) for automated illumination restoration. It takes the decoupled 3D illumination map $\hat{L}$ as input, predicts an content-adaptive gamma correction parameter g , and generates the enhanced illumination map $L_{\text{enh}} = (\hat{L})^g$. It is worth noting that while physical light transport is linear, applying Gamma correction to the illumination component is a well-established practice in Retinex-based low-light enhancement. This operation functions as a non-linear tone mapping operator, approximating the dynamic range compression of the Human Visual System. By adjusting the gamma value, we can effectively boost visibility in under-exposed regions while preventing saturation in well-lit areas. The final normal-light image is synthesized as $S_{\text{enh}} = \hat{R} \cdot L_{\text{enh}}$. The training process of EnhancementNet is completely unsupervised, guided by a composite loss L_{enhance} imposed on the final output S_{enh} :

$$L_{\text{enhance}} = \lambda_{\text{spa}} \cdot L_{\text{spa}} + \lambda_{\text{exp}} \cdot L_{\text{exp}} + \lambda_{\text{col}} \cdot L_{\text{col}}, \quad (7)$$

where L_{spa} preserves textures by penalizing the L1 norm of gradient differences between the enhanced image S_{enh} and the original image S :

$$L_{\text{spa}} = \|\nabla S_{\text{enh}} - \nabla S\|_1, \quad (8)$$

L_{exp} drives the average brightness Y_k of M local image patches to approach the target exposure value E to prevent overexposure or underexposure:

$$L_{\text{exp}} = \frac{1}{M} \cdot \sum |Y_k - E|, \quad (9)$$

and L_{col} corrects color bias by minimizing the variance between the mean values of RGB channels based on the gray-world assumption [26]:

$$L_{\text{col}} = \sum_{p,q \in \{R,G,B\}} (S_{\text{mean}_p} - S_{\text{mean}_q})^2. \quad (10)$$

Here, λ_{spa} , λ_{exp} , and λ_{col} are set to 0.1, 1, and 0.5, respectively. To stabilize training, these unsupervised losses are introduced after an initial 3000 steps of optimization using only the decomposition loss, allowing the network to learn robust illumination restoration.

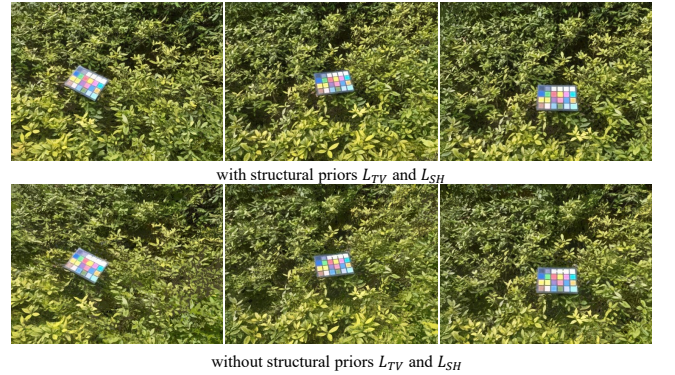


Fig. 5. Ablation study on the structural regularization priors. Removing L_{TV} and L_{SH} leads to a noticeable increase in noise and artifacts in the image.



Fig. 6. Ablation study on the unsupervised enhancement loss L_{enhance} . Removing L_{spa} causes blurry textures, the L_{exp} leads to underexposure, and the L_{col} results in a color cast.

III. EXPERIMENTS

A. Experimental Setup & Dataset

Our method is implemented based on PyTorch [27] and the open-source GS-Splat [28] toolbox. The entire model is trained on a single NVIDIA RTX 4090 GPU, with 10,000 training steps per scene, taking approximately 10 minutes. It is also worth noting that the introduced MLP incurs low computational cost, achieving an average rendering speed of 96 FPS on an RTX 4090 GPU, thus maintaining the real-time rendering capability of 3DGS.

To comprehensively evaluate the effectiveness and robustness of RetinexGS, following previous works ([16]–[18]), we conducted experiments on the LOM dataset [17], which serves

as the de facto standard benchmark for low-light novel view synthesis. This real-world dataset is specifically designed for this task, containing indoor and outdoor scenes, and each scene provides paired low-light and normal-light images captured under the same camera pose. Following its standard division, we randomly selected 3 to 5 views as the test set in each scene, with the rest serving as the training set.

B. Comparative Results

To evaluate the performance of the proposed method, we selected PSNR, SSIM, and LPIPS as core metrics. The comparison objects include three categories: first, the original 3DGS [2] model, which is used to verify the performance upper limit of the basic model in low-light scenarios; second, two schemes that combine 3DGS with SOTA 2D image and video enhancement techniques [11]–[14]—one using enhancement methods to preprocess training images and the other applying enhancement methods to postprocess output images during the rendering stage; third, three representative and advanced peer methods, namely Aleth-nerf [17], LLGS [16], and Luminance-GS [18].

Quantitative comparison. Table I presents the results of quantitative comparisons on the LOM dataset. Our RetinexGS achieves the best average performance across all three metrics, and particularly achieves the best scores across the board in terms of LPIPS, a metric for evaluating perceptual quality. This leading advantage highlights the superiority of our hybrid supervision paradigm: unlike 2D preprocessing schemes that tend to introduce view-inconsistent artifacts, our method ensures 3D consistency; in contrast to other 3D decomposition methods that rely solely on unsupervised weak priors, RetinexGS incorporates 2D strong priors for guidance, resulting in visually plausible decomposition outcomes and thus synthesizing images with higher perceptual quality.

Qualitative comparison. To intuitively assess the performance of our method, qualitative comparison results between RetinexGS and several state-of-the-art methods are presented in Figure 4. From the visual comparisons, it can be observed that our method exhibits advantages in detail restoration, artifact suppression, and color fidelity. In contrast to the per-Gaussian modulation schemes [16], [18], where discrete coupling leads to overfitting and artifacts, our method achieves a cleaner disentanglement of illumination and reflectance, thereby demonstrating superior spatial consistency and robustness in low-light scenarios characterized by severe noise.

C. Ablation Results

Impact of Structural Regularization. As shown in Figure 5 and Table II, removing the structural regularization priors L_{TV} and L_{SH} leads to a decline in metrics and significant noise in the rendered image. This confirms that these constraints are essential for decoupling view-dependent effects and maintaining geometric consistency.

Impact of Enhancement Loss. Regarding the enhancement loss $L_{enhance}$, Table II quantifies the critical role of each term. Notably, removing the exposure loss L_{exp} causes a catastrophic

TABLE II
THE QUANTITATIVE ABLATION RESULTS ON THE LOM DATASET (PSNR $\uparrow$, SSIM $\uparrow$, LPIPS $\downarrow$). BEST RESULTS ARE BOLDED.

Settings	PSNR $\uparrow$	SSIM $\uparrow$	LPIPS $\downarrow$
w/o L_{TV}	20.45	0.755	0.285
w/o L_{SH}	20.92	0.792	0.245
w/o L_{spa}	20.15	0.730	0.255
w/o L_{exp}	12.45	0.467	0.523
w/o L_{col}	20.11	0.805	0.275
w/ RUAS	21.76	0.829	0.217
Full model	21.85	0.827	0.211

drop in all metrics, as the model fails to restore brightness and outputs under-exposed images. Removing the spatial loss L_{spa} leads to the loss of texture details. Finally, removing the color loss L_{col} causes color casts, as shown in Figure 6.



Fig. 7. Ablation study on the choice of 2D decomposer. The comparable visual quality demonstrates the robustness of our hybrid supervision strategy to different Retinex decomposition methods.



Fig. 8. Application to real-world scenarios. Top: Original low-light inputs characterized by severe visibility degradation. Bottom: Reconstruction results by RetinexGS. Our method demonstrates robust generalization in uncontrolled environments, successfully generating clear geometry and coherent colors from data with almost zero visibility.

Impact of 2D Decomposer. To investigate the robustness of our framework regarding the choice of 2D priors, we conducted an experiment using RUAS [22], another representative Retinex decomposition method, as a replacement for Lr3m [25]. As shown in Table II and Figure 7, the performance with RUAS is comparable to, albeit slightly inferior to, that of Lr3m. This indicates that RetinexGS is not heavily dependent on a specific 2D decomposer, demonstrating the robustness of our hybrid supervision strategy. Furthermore, removing the 2D priors entirely leads to failed reconstruction.

D. Generalization on Real-World Scenes

While the LOM dataset serves as the primary benchmark, real-world low-light environments often present more complex

challenges. As a supplementary qualitative evaluation, we applied RetinexGS to a few selected 'in-the-wild' scenarios. As illustrated in Figure 8, our framework successfully reconstructs visually plausible 3D scenes from these challenging inputs. These results preliminarily demonstrate the potential of our strategy to handle uncontrolled environments beyond standard datasets.

IV. CONCLUSIONS

This paper presents RetinexGS, a novel framework addressing the challenge of low-light novel view synthesis via a hybrid supervision paradigm. To mitigate the instability of pure unsupervised 3D decomposition, our method leverages strong priors from a mature 2D Retinex model, guiding the decoupling of scenes into a view-consistent reflectance field and a structure-aware illumination field. Furthermore, a dedicated enhancement network is introduced to automatically restore the scenes to normal brightness. Extensive experiments demonstrate that our approach achieves performance that is competitive with, and often superior to, the current state-of-the-art. However, our performance relies on the quality of 2D priors, particularly in scenarios with extreme degradation. Future work aims to integrate the decomposition directly into the 3D optimization process to bypass the offline preprocessing stage.

REFERENCES

- [1] Ben Mildenhall, Pratul P Srinivasan, Matthew Tancik, Jonathan T Barron, Ravi Ramamoorthi, and Ren Ng, "Nerf: Representing scenes as neural radiance fields for view synthesis," *Communications of the ACM*, vol. 65, no. 1, pp. 99–106, 2021.
- [2] Bernhard Kerbl, Georgios Kopanas, Thomas Leimkühler, and George Drettakis, "3d gaussian splatting for real-time radiance field rendering," *ACM Trans. Graph.*, vol. 42, no. 4, pp. 139–1, 2023.
- [3] Zehao Yu, Anpei Chen, Binbin Huang, Torsten Sattler, and Andreas Geiger, "Mip-splatting: Alias-free 3d gaussian splatting," in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2024, pp. 19447–19456.
- [4] Tao Lu, Mulin Yu, Linning Xu, Yuanbo Xiangli, Limin Wang, Dahua Lin, and Bo Dai, "Scaffold-gs: Structured 3d gaussians for view-adaptive rendering," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2024, pp. 20654–20664.
- [5] Binbin Huang, Zehao Yu, Anpei Chen, Andreas Geiger, and Shenghua Gao, "2d gaussian splatting for geometrically accurate radiance fields," in *ACM SIGGRAPH 2024 conference papers*, 2024, pp. 1–11.
- [6] Lukas Radl, Michael Steiner, Mathias Parger, Alexander Weinrauch, Bernhard Kerbl, and Markus Steinberger, "Stopthepop: Sorted gaussian splatting for view-consistent real-time rendering," *ACM Transactions on Graphics (TOG)*, vol. 43, no. 4, pp. 1–17, 2024.
- [7] Kai Cheng, Xiaoxiao Long, Kaizhi Yang, Yao Yao, Wei Yin, Yuexin Ma, Wenping Wang, and Xuejin Chen, "Gaussianpro: 3d gaussian splatting with progressive propagation," in *Forty-first International Conference on Machine Learning*, 2024.
- [8] Jiahui Zhang, Fangneng Zhan, Muyu Xu, Shijian Lu, and Eric Xing, "Fregs: 3d gaussian splatting with progressive frequency regularization," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2024, pp. 21424–21433.
- [9] Yingwenqi Jiang, Jiadong Tu, Yuan Liu, Xifeng Gao, Xiaoxiao Long, Wenping Wang, and Yuexin Ma, "Gaussianshader: 3d gaussian splatting with shading functions for reflective surfaces," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2024, pp. 5322–5332.
- [10] L Franke, D Rückert, L Fink, and M Stamminger, "Trips: Trilinear point splatting for real-time radiance field rendering. arxiv abs/2401.06003 (2024)," .
- [11] Chen Wei, Wenjing Wang, Wenhan Yang, and Jiaying Liu, "Deep retinex decomposition for low-light enhancement," *arXiv preprint arXiv:1808.04560*, 2018.
- [12] Chunle Guo, Chongyi Li, Jichang Guo, Chen Change Loy, Junhui Hou, Sam Kwong, and Runmin Cong, "Zero-reference deep curve estimation for low-light image enhancement," in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2020, pp. 1780–1789.
- [13] Fan Zhang, Yu Li, Shaodi You, and Ying Fu, "Learning temporal consistency for low light video enhancement from single images," in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2021, pp. 4967–4976.
- [14] Long Ma, Tengyu Ma, Risheng Liu, Xin Fan, and Zhongxuan Luo, "Toward fast, flexible, and robust low-light image enhancement," in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2022, pp. 5637–5646.
- [15] Haoyuan Wang, Xiaogang Xu, Ke Xu, and Rynson WH Lau, "Lighting up nerf via unsupervised decomposition and enhancement," in *Proceedings of the IEEE/CVF international conference on computer vision*, 2023, pp. 12632–12641.
- [16] Jianwen Gan, Wenxin Li, Bo Zheng, Chengliang Wang, and Yingbo Wu, "Llgs: Illuminating gaussian splatting via absorptance modulation," in *ICASSP 2025-2025 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 2025, pp. 1–5.
- [17] Ziteng Cui, Lin Gu, Xiao Sun, Xianzheng Ma, Yu Qiao, and Tatsuya Harada, "Aleth-nerf: Illumination adaptive nerf with concealing field assumption," in *Proceedings of the AAAI conference on artificial intelligence*, 2024, vol. 38, pp. 1435–1444.
- [18] Ziteng Cui, Xuangeng Chu, and Tatsuya Harada, "Luminance-gs: Adapting 3d gaussian splatting to challenging lighting conditions with view-adaptive curve adjustment," in *Proceedings of the Computer Vision and Pattern Recognition Conference*, 2025, pp. 26472–26482.
- [19] Edwin H Land, "The retinex theory of color vision," *Scientific american*, vol. 237, no. 6, pp. 108–129, 1977.
- [20] Daniel J Jobson, Zia-ur Rahman, and Glenn A Woodell, "A multiscale retinex for bridging the gap between color images and the human observation of scenes," *IEEE Transactions on Image processing*, vol. 6, no. 7, pp. 965–976, 1997.
- [21] Daniel J Jobson, Zia-ur Rahman, and Glenn A Woodell, "Properties and performance of a center/surround retinex," *IEEE transactions on image processing*, vol. 6, no. 3, pp. 451–462, 1997.
- [22] Risheng Liu, Long Ma, Jiaao Zhang, Xin Fan, and Zhongxuan Luo, "Retinex-inspired unrolling with cooperative prior architecture search for low-light image enhancement," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2021, pp. 10561–10570.
- [23] Yuanhao Cai, Hao Bian, Jing Lin, Haoqian Wang, Radu Timofte, and Yulun Zhang, "Retinexformer: One-stage retinex-based transformer for low-light image enhancement," in *Proceedings of the IEEE/CVF international conference on computer vision*, 2023, pp. 12504–12513.
- [24] Ben Fei, Zhaoyang Lyu, Liang Pan, Junzhe Zhang, Weidong Yang, Tianyue Luo, Bo Zhang, and Bo Dai, "Generative diffusion prior for unified image restoration and enhancement," in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2023, pp. 9935–9946.
- [25] Xutong Ren, Wenhan Yang, Wen-Huang Cheng, and Jiaying Liu, "Lr3m: Robust low-light enhancement via low-rank regularized retinex model," *IEEE Transactions on Image Processing*, vol. 29, pp. 5862–5876, 2020.
- [26] Edmund Y Lam, "Combining gray world and retinex theory for automatic white balance in digital photography," in *Proceedings of the Ninth International Symposium on Consumer Electronics, 2005.(ISCE 2005)*. IEEE, 2005, pp. 134–139.
- [27] Adam Paszke, Sam Gross, Francisco Massa, Adam Lerer, James Bradbury, Gregory Chanan, Trevor Killeen, Zeming Lin, Natalia Gimelshein, Luca Antiga, et al., "Pytorch: An imperative style, high-performance deep learning library," *Advances in neural information processing systems*, vol. 32, 2019.
- [28] Vickie Ye, Ruilong Li, Justin Kerr, Matias Turkulainen, Brent Yi, Zhuoyang Pan, Otto Seiskari, Jianbo Ye, Jeffrey Hu, Matthew Tancik, et al., "gsplat: An open-source library for gaussian splatting," *Journal of Machine Learning Research*, vol. 26, no. 34, pp. 1–17, 2025.