

InstantHDR: Single-forward Gaussian Splatting Initialization for HDR 3D Reconstruction

Dingqiang Ye^{*1}, Jiacong Xu^{*1}, Jianglu Ping¹,
Yuxiang Guo¹, Chao Fan², and Vishal M. Patel^{†,1}

¹ Johns Hopkins University, USA

² Shenzhen University, China

{dye6, jxu155, jping1, yguo87}@jhu.edu,
chaofan996@szu.edu.cn, vpatel36@jhu.edu

Abstract. High dynamic range (HDR) novel view synthesis (NVS) aims to reconstruct HDR scenes from multi-exposure low dynamic range (LDR) images. Existing HDR pipelines heavily rely on known camera poses, well-initialized dense point clouds, and time-consuming per-scene optimization. Current feed-forward alternatives overlook the HDR problem by assuming exposure-invariant appearance. To bridge this gap, we propose InstantHDR, a feed-forward network that initializes 3D HDR scenes from uncalibrated multi-exposure LDR collections in a fast single forward pass. Specifically, we design a geometry-guided appearance modeling for multi-exposure fusion, and a meta-network for generalizable scene-specific tone mapping. Due to the lack of HDR scene data, we build a pre-training dataset, called HDR-Pretrain, for generalizable feed-forward HDR models, featuring 168 Blender-rendered scenes, diverse lighting types, and multiple camera response functions. Comprehensive experiments show that our InstantHDR delivers a single-forward HDR initialization at $\sim 700\times$ the speed of SoTA optimization-based methods, and reaches comparable quality in real settings after lightweight post-optimization while remaining $\sim 20\times$ faster. All code, models, and datasets: <https://github.com/Bugjudger/InstantHDR>.

Keywords: High Dynamic Range · 3D Gaussian Splatting · Novel-View Synthesis · Feed-Forward Models

1 Introduction

High dynamic range (HDR) novel view synthesis (NVS) aims to reconstruct HDR scenes from multi-view low dynamic range (LDR) images captured at varying exposures. Unlike typical low dynamic range (from 0 to 255) imaging, which often suffers from detail loss in extreme lighting and color distortion due to sensor limitations, HDR captures a broader spectrum of luminance (from 0 to $+\infty$). Through integrating advanced frameworks like NeRF [55] or 3D Gaussian Splatting [28] with a tone mapper to model the camera response function (CRF), HDR

^{*} co-first authors; [†]corresponding author

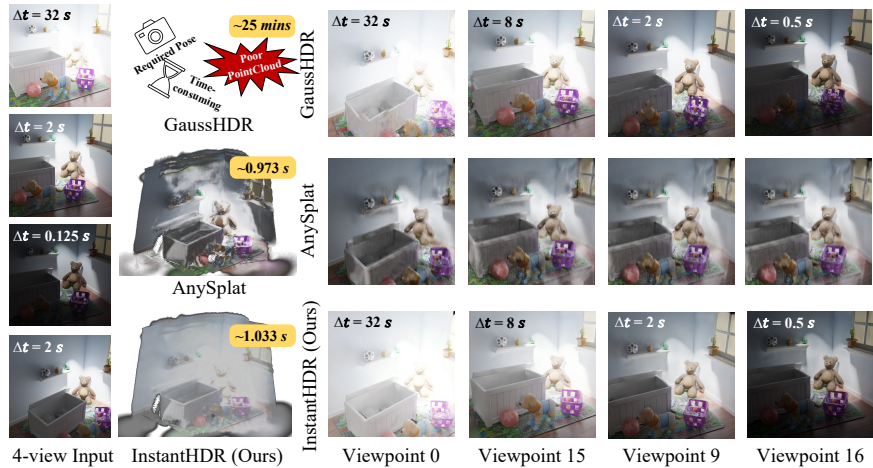


Fig. 1: Comparisons of reconstruction time (yellow boxes), scenes (left) and rendered views (right) between the GaussHDR [41] (top), original AnySplat [21] (middle) and our InstantHDR (bottom). (i) GaussHDR [41] spends expensive 25 mins and produces tearing artifacts, as its initial point clouds collapse under the sparse-view inputs. (ii) AnySplat [21] naively fuses multi-exposure inputs, causing ghosting artifacts and lacking exposure control. (iii) Our InstantHDR initializes 3D-consistent HDR scenes in few seconds and renders clean LDR images with controllable exposure time.

task enables the re-rendering of photo-realistic novel views with controllable exposure. Its ability to faithfully represent real-world light and shadow makes it indispensable for multiple applications such as autonomous driving [16, 67, 71, 87], digital humans [15, 43, 98, 99], and immersive image editing [32, 45, 66, 92].

Existing HDR NVS methods [3, 17, 41] are predominantly optimization-based. However, this paradigm is costly and generalizes poorly: it relies on precisely calibrated camera poses and dense multi-view inputs for SfM-based point cloud initialization, followed by time-consuming per-scene densification and optimization. For example, as shown in Fig. 1, GaussHDR [41] struggles to reconstruct scenes from only four LDR views, since exposure-induced appearance inconsistencies degrade the reliability of SfM-based point cloud initialization in sparse-view settings. Furthermore, their heavy computational overhead and strong data dependency limit practical deployment in real-time scenarios.

Recently, 3D feed-forward models [21, 73] have revolutionized scene reconstruction by inferring geometry in seconds, achieving impressive speed and better generalization than optimization-based methods. Integrating this paradigm into the HDR NVS tasks could boost model generalizability and inference speed. However, directly applying the original feed-forward models to HDR reconstruction may encounter the following issues. (a) **Exposure-induced Appearance Inconsistency:** Naive fusion leads to severe ghosting — as shown in Fig. 1. The same white wall appears bright at $\Delta = 32s$ but nearly black at $\Delta = 0.125s$, caus-

ing visible artifacts in AnySplat [21]. (b) **Pixel-level Geometric Alignment:** Establishing accurate pixel-level geometric correspondences remains a non-trivial task under large brightness variations. (c) **Camera Response Functions Inconsistency:** In real world, different camera and software apply distinct color transformations (*e.g.*, AgX, Filmic), making it difficult to learn a unified tone mapping operator. (d) **HDR Data Scarcity:** Current publicly available HDR datasets [17, 25] are insufficient (as shown in Tab. 1) to support the robust large-scale pre-training required for feed-forward models.

To address these challenges, we propose **InstantHDR**, a novel feed-forward 3D reconstruction framework for HDR novel view synthesis. InstantHDR comprises two key components. First, a *geometry-guided appearance modeling module* normalizes multi-exposure LDR inputs into a unified exposure space and utilizes geo-attention from the geometry encoder to fuse patch-level irradiance features. Fine-grained texture details are further recovered by incorporating Difference of Gaussians (DoG) high-frequency cues, lifting the representation to pixel-level irradiance before predicting HDR 3D Gaussians. Second, a *MetaNet* takes the predicted HDR Gaussians, LDR images, and exposure times as input to estimate scene-specific tonemapper parameters, enabling generalizable HDR-to-LDR rendering with controllable exposure time Δt . This fast single forward pass yields a strong HDR initialization; a lightweight post-optimization stage then refines it to quality comparable to SoTA, while together remaining orders of magnitude faster than existing fully optimization-based pipelines.

Our contributions can be summarized as follows:

(i) We propose InstantHDR, the first feed-forward HDR novel view synthesis method. It features a geometry-guided appearance module for exposure-robust multi-view fusion and a meta-network for generalizable tone mapping, enabling 3D HDR initialization from uncalibrated multi-exposure LDRs in seconds.

(ii) We build HDR-Pretrain, a large-scale dataset of 168 synthetic indoor scenes to support feed-forward HDR pretraining. Experiments show that InstantHDR initializes HDR scenes $\sim 700\times$ faster than SoTA, reaching competitive quality after lightweight post-optimization while staying $\sim 20\times$ faster.

2 Related Works

High Dynamic Range Imaging. HDR imaging has traditionally merged multiple LDR exposures from a fixed viewpoint [53] or recovered the camera response function from bracketed LDR sequences [8, 77, 78], but these methods suffer from ghosting under scene motion. Subsequent works [12, 18, 26, 70, 84] mitigate this via optical-flow-based motion compensation prior to fusion, while learning-based approaches directly learn LDR-to-HDR mappings using CNNs [9, 29, 31, 48] and Transformers [4, 23, 49, 65, 90, 91]. Others reconstruct HDR from a single LDR image via handcrafted priors in an unsupervised or self-supervised manner [10, 38, 57, 60, 83, 97]. A parallel line operates in the RAW domain, exploiting its higher bit-depth and linear radiometric response—from a single RAW image [101], by jointly denoising and fusing multi-exposure RAW

frames [19, 34, 44], or by reconstructing HDR video from alternating-exposure RAW sequences [62, 81], often coupled with the in-camera ISP [36]. However, these methods lack the 3D understanding needed for HDR novel views.

Gaussian Splatting. 3D Gaussian Splatting (3DGS) [28] represents scenes as collections of anisotropic Gaussian primitives, enabling real-time rendering through efficient rasterization—offering a significant speed advantage over NeRF-based volumetric ray-marching [17, 25, 54]. This efficiency has driven its adoption across diverse tasks including dynamic scenes [51, 79, 86], SLAM [27, 52, 82, 93], inverse rendering [22, 40, 80], digital humans [14, 33, 46], 3D generation [39, 68, 89], and medical imaging [2, 94]. Nevertheless, current HDR extensions [1, 3, 6, 11, 24, 37, 41, 42, 63, 100] of 3DGS remain predominantly optimization-based, resulting in expensive per-scene reconstruction times. Our work aims to fill this gap.

Feed-forward 3D Reconstruction. Recent advances pursue end-to-end 3D reconstruction directly from unposed images. Pioneering works such as DUST3R [76] and MAST3R [35] replace multi-stage pipelines with a unified model that jointly estimates depth and performs dense scene fusion, and subsequent approaches [47, 56, 69, 72, 73, 75, 85] cascade transformer blocks to recover camera poses, point trajectories, and scene geometry in a single forward pass. A parallel line [5, 13, 20, 64, 74, 88, 96] targets novel view synthesis from unposed sparse-view images. Compared to optimization-based ones, these models offer remarkable speed, strong generalization, and minimal data requirements, yet their potential for HDR reconstruction remains largely unexplored. We explore this promising direction.

3 Method

Given uncalibrated multi-view LDR images captured at varying exposures, InstantHDR initializes HDR 3D Gaussians and renders novel views at any target exposure (Fig. 2). We first formalize the problem in Sec. 3.1, then detail the pipeline design in Sec. 3.2, including a Geo-guided Appearance Modeling module (Sec. 3.3) and a 3D HDR-to-2D LDR Mapping module (Sec. 3.4), and finally present our training strategies in Sec. 3.5.

3.1 Problem Setup

Inputs. Consider V *uncalibrated* views of a single 3D scene, given as context images $\{I_v\}_{v=1}^V$, $I_v \in \mathbb{R}^{H \times W \times 3}$, each captured at a known exposure time Δt_v . For convenience, we work in log-exposure space and define $\ell_v = \log_2 \Delta t_v$.

Outputs. InstantHDR jointly reconstructs the scene geometry and HDR appearance by predicting:

(a) *HDR 3D Gaussians.* A collection of G anisotropic 3D Gaussians

$$\{(\boldsymbol{\mu}_g, \sigma_g, \mathbf{r}_g, \mathbf{s}_g, \mathbf{c}_g^h)\}_{g=1}^G, \quad (1)$$

where each Gaussian is parameterized by a center position $\boldsymbol{\mu} \in \mathbb{R}^3$, an opacity $\sigma \in \mathbb{R}^+$, an orientation quaternion $\mathbf{r} \in \mathbb{R}^4$, an anisotropic scale $\mathbf{s} \in \mathbb{R}^3$, and

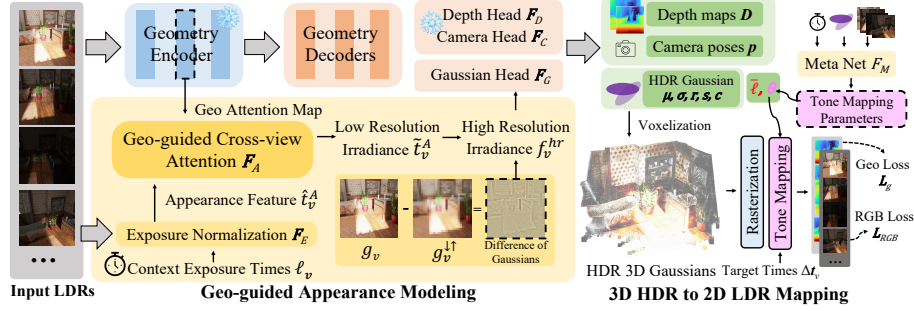


Fig. 2: Overview of InstantHDR. Given multi-exposure LDR images, the frozen geometry branch estimates depth and camera poses, while the appearance branch normalizes exposures (F_E), fuses cross-view irradiance via geometry-guided attention (F_A), and recovers pixel-level details via DoG upsampling. The Gaussian head F_G combines both branches to produce HDR 3D Gaussians. The Meta Net F_M predicts tone-mapping parameters for rendering LDR images at controllable exposures.

an HDR color embedding $c^h \in \mathbb{R}^{3 \times (k+1)^2}$ parameterized as degree- k spherical-harmonic (SH) coefficients that encode *log-radiance*, following [3, 41].

(b) *Camera parameters.* Per-view parameters $\{p_v \in \mathbb{R}^9\}_{v=1}^V$, where p_v comprises a focal length, a 3-DoF rotation (axis-angle), a 3-DoF translation, and a 2-DoF principal-point offset.

(c) *Scene-level attributes.* A mid-exposure anchor $\bar{\ell} = \frac{1}{2}(\max_v \ell_v + \min_v \ell_v)$, serving as the reference exposure level, and the parameters θ of a lightweight tonemapper that approximates the scene-specific camera response function (CRF).

Overall mapping. Formally, our model implements:

$$f_{\Theta}: \{I_v, \ell_v\}_{v=1}^V \mapsto \left\{ (\mu_g, \sigma_g, r_g, s_g, c_g^h) \right\}_{g=1}^G \cup \{p_v\}_{v=1}^V \cup \{\bar{\ell}, \theta\}. \quad (2)$$

3.2 Pipeline Overview

As illustrated in Fig. 2, our pipeline consists of two branches: a *geometry branch* that estimates scene structure from multi-view images, and an *appearance branch* that reconstructs HDR irradiance from exposure-inconsistent inputs. Given V uncalibrated multi-exposure LDR images, the geometry branch encodes them into high-dimensional features via a pretrained transformer and decodes depth maps and camera poses. The appearance branch—our core contribution—uses a *Geo-guided Appearance Modeling* module that leverages geometric correspondences from the geometry branch to fuse multi-exposure information into a coherent HDR representation. The outputs of both branches are combined by a Gaussian head to produce HDR 3D Gaussians, which are then converted to LDR via a learned tonemapper (Sec. 3.4).

Geometry Branch. The geometry branch provides the structural foundation for our pipeline. Following VGGT [73] and AnySplat [21], we adopt a pretrained alternating-attention transformer as the geometry backbone. Each image I_v is patchified into $N = \frac{HW}{p^2}$ tokens of dimension d using DINOv2 [58], where $p=14$ and $d=1024$. To each token sequence $\mathbf{t}_v^G \in \mathbb{R}^{N \times d}$, we prepend a learnable camera token $\mathbf{t}_v^{\text{cam}} \in \mathbb{R}^{1 \times d}$ and four register tokens $\mathbf{t}_v^R \in \mathbb{R}^{4 \times d}$. The combined tokens from all V views are processed by an L -layer alternating-attention transformer, where each layer applies frame-wise self-attention followed by global cross-view attention. Dedicated decoder heads for camera poses p_v and depth maps D_v are *frozen* together with the geometry encoder throughout training.

Gaussian Head. The Gaussian head combines information from both branches to predict HDR-aware Gaussian attributes. It takes the geometry tokens $\mathbf{t}_v^G$ and the high-resolution irradiance features $\mathbf{f}_v^{\text{hr}}$ produced by the Geo-guided Appearance Modeling module (Sec. 3.3) as input. The Gaussian head remains *trainable*, enabling it to output HDR-aware Gaussian attributes $\{\sigma_g, \mathbf{r}_g, \mathbf{s}_g, \mathbf{c}_g^h\}$.

3.3 Geo-guided Appearance Modeling

The frozen geometry branch provides reliable structure but cannot handle exposure-induced appearance inconsistency. We introduce the Geo-guided Appearance Modeling module, which mitigates this problem through three stages: (1) *Exposure Normalization* aligns inputs to a common reference level, (2) *Geo-guided Cross-view Attention* fuses irradiance features using geometric correspondences from the frozen backbone, and (3) *High-Resolution Upsampling* recovers pixel-level textures via high-frequency cues. The output features are decoded into log-radiance SH colors $\mathbf{c}_g^h$ for each Gaussian.

Exposure Normalization F_E . To fuse multi-exposure views, we first remove the exposure-induced brightness variation by normalizing all appearance features to a shared reference level. We define the relative log-exposure of each view as $\tilde{\ell}_v = \ell_v - \bar{\ell}$, where $\bar{\ell}$ is the mid-exposure anchor from Eq. (2), and encode $\tilde{\ell}_v$ into a d -dimensional embedding $\mathbf{e}_v$ via sinusoidal positional encoding [55]. Meanwhile, we extract per-view appearance tokens $\mathbf{t}_v^A \in \mathbb{R}^{N \times d}$ from each LDR image using a separate patch encoder with the same patch size p and dimension d as the geometry backbone. A FiLM layer [59] then predicts per-view affine parameters:

$$(\gamma_v, \beta_v) = \text{FiLM}(\mathbf{e}_v, \bar{\mathbf{a}}_v, \bar{\mathbf{a}}), \quad (3)$$

where $\bar{\mathbf{a}}_v = \frac{1}{N} \sum_{n=1}^N \mathbf{t}_{v,n}^A$ is the per-view feature mean and $\bar{\mathbf{a}} = \frac{1}{V} \sum_{v=1}^V \bar{\mathbf{a}}_v$ is the global scene summary. The appearance tokens are then modulated as:

$$\hat{\mathbf{t}}_v^A = \mathbf{t}_v^A \odot (1 + \gamma_v) + \beta_v, \quad (4)$$

aligning all views to a shared irradiance level before cross-view fusion. FiLM is not meant to invert the full non-linear CRF, but to coarsely align exposure distributions into a comparable range before fusion. The remaining spatially varying

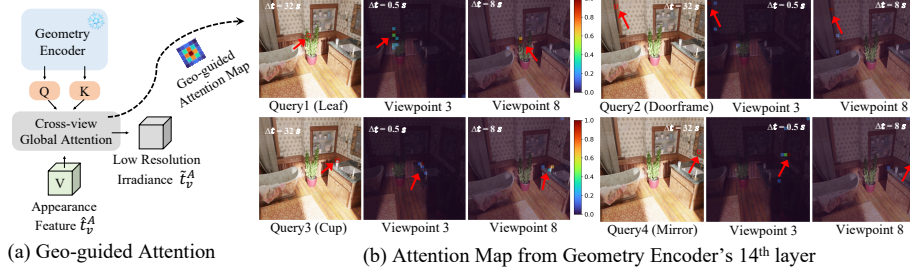


Fig. 3: Geo-guided Cross-view Attention. (a) The module reuses Q, K from the 14th frozen geometry encoder layer to guide appearance fusion. (b) Attention maps visualization shows that it naturally and accurately matches query patches (red box) across views under large viewpoint and extreme exposure variations (Δt : 0.5–32s).

effects, which a global FiLM cannot capture, are handled by the subsequent patch-level Geo-guided Attention. Note that, unlike methods that first linearize inputs via inverse gamma correction [97], our model operates directly on camera-output LDR images (gamma-encoded), since our training objective reconstructs multi-exposure LDR images rather than linear HDR radiance (Sec. 3.5).

Geo-guided Cross-view Attention F_A . Different exposures capture complementary information: bright exposures reveal shadows while dark ones preserve highlights. Fusing them requires cross-view correspondences, which are challenging under large viewpoint and exposure changes. Interestingly, we observe that *the global attention maps in the frozen geometry encoder already encode reliable cross-view geometric correspondences, greatly benefiting our appearance modeling*. Therefore, we reuse these attention maps to guide appearance fusion. As shown in Fig. 3 (b), diverse elements such as leaves, cups, doorframes, and mirrors are accurately matched across views despite extreme exposure variations.

$$\tilde{t}_v^A = \text{softmax}\left(\frac{QK^\top}{\sqrt{d}}\right) \hat{t}_v^A. \quad (5)$$

High-Resolution Upsampling F_U . The irradiance features $\tilde{t}_v^A$ operate at patch resolution $(\frac{H}{p} \times \frac{W}{p})$, losing high-frequency details critical for realistic appearance. To recover pixel-level textures while preserving the fused low-frequency irradiance, we adopt a Difference-of-Gaussians (DoG) guided upsampling strategy. We first encode each full-resolution LDR image into a feature map $\mathbf{g}_v \in \mathbb{R}^{d' \times H \times W}$ via a shallow CNN. A low-passed version $\mathbf{g}_v^{\downarrow\uparrow}$ is obtained by downsampling and bilinearly upsampling $\mathbf{g}_v$, yielding the high-frequency residual $\mathbf{g}_v - \mathbf{g}_v^{\downarrow\uparrow}$. This residual is then added to the bilinearly upsampled irradiance features to produce pixel-level irradiance features:

$$\mathbf{f}_v^{\text{hr}} = \text{Conv}(\text{Up}(\tilde{t}_v^A) + \text{Conv}(\mathbf{g}_v - \mathbf{g}_v^{\downarrow\uparrow})) \in \mathbb{R}^{d \times H \times W}, \quad (6)$$

combining multi-view irradiance consensus with per-image structural detail.

HDR 3D Gaussian Prediction. We now merge the geometry and appearance branches to produce HDR 3D Gaussians. A DPT decoder [61] upsamples the geometry tokens $\mathbf{t}_v^G$ to pixel resolution, and the result is added to the irradiance features $\mathbf{f}_v^{\text{hr}}$. A lightweight CNN then regresses per-Gaussian opacity, orientation, scale, and log-radiance SH color:

$$\{\sigma_g, \mathbf{r}_g, \mathbf{s}_g, \mathbf{c}_g^h\} = F_G(\text{DPT}(\mathbf{t}_v^G) + \mathbf{f}_v^{\text{hr}}). \quad (7)$$

The Gaussian centers $\boldsymbol{\mu}_g$ are obtained by back-projecting the predicted depth maps D_v through the estimated camera poses p_v . The Gaussians are then voxelized [21] to reduce primitive count for efficient splatting.

3.4 3D HDR-to-2D LDR Mapping

Given HDR 3D Gaussians, rendering a photorealistic LDR view requires recovering the Camera Response Function (CRF) that maps scene irradiance to observed pixel values. Unlike optimization-based methods that overfit a per-scene MLP tonemapper [3, 41], our method seeks a *generalizable* tonemapper that adapts to different cameras without per-scene optimization. We achieve this goal via a **Meta Net** F_M that predicts the parameters of a lightweight tonemapper from scene context.

Tone Mapping Formulation. We convert HDR radiance to LDR in two steps. First, Gaussian splatting rasterizes the linear radiance into a per-pixel HDR image at view v :

$$\mathbf{H}_v = \mathcal{R}(\exp(\mathbf{c}_g^h), \sigma_g, \mathbf{r}_g, \mathbf{s}_g, \boldsymbol{\mu}_g; p_v) \in \mathbb{R}^{H \times W \times 3}, \quad (8)$$

where $\mathcal{R}(\cdot)$ denotes the differentiable rasterization with alpha-weighted blending [28] in linear radiance space, p_v is the camera pose, and the $\exp$ converts log-radiance SH colors $\mathbf{c}_g^h$ back to linear radiance before blending.

Second, following the log-domain CRF model [8], a learned tonemapper g_θ maps the log-irradiance to $[0, 1]$ LDR values:

$$\mathbf{L}_v(\ell) = g_\theta(\log \mathbf{H}_v + (\ell - \bar{\ell}) \cdot \log 2), \quad (9)$$

where $\mathbf{L}_v(\ell) \in \mathbb{R}^{H \times W \times 3}$ is the rendered LDR image at view v under target log-exposure ℓ , $\bar{\ell}$ is the mid-exposure anchor from Eq. (2), and the term $(\ell - \bar{\ell}) \cdot \log 2$ adjusts the irradiance to the desired exposure level. Here g_θ is a two-layer MLP with hidden dimension h (input: 3 $\rightarrow$ hidden: h with ReLU $\rightarrow$ output: 3 with sigmoid), acting as a learnable inverse CRF. Rather than learning a fixed θ , we predict *scene-specific* parameters via the Meta Net, enabling adaptation to different cameras and tone curves without per-scene optimization.

Meta Net F_M . The Meta Net infers scene-specific tonemapper parameters θ in a single forward pass, enabling g_θ to reproduce the original camera’s tone curve

without per-scene optimization. For brevity, we denote the full set of pixel-level Gaussians (before voxelization) as:

$$\mathcal{G} = \{(\boldsymbol{\mu}_g, \sigma_g, \mathbf{r}_g, \mathbf{s}_g, \mathbf{c}_g^h)\}_{g=1}^G, \quad G = V \times H \times W. \quad (10)$$

The Meta Net takes three inputs: (i) the full-resolution LDR features $\mathbf{g}_v$ from the upsampling CNN (Eq. (6)), (ii) the per-view exposure embeddings $\{\mathbf{e}_v\}_{v=1}^V$, and (iii) the predicted HDR Gaussians $\mathcal{G}$. These are concatenated and compressed by a strided convolutional encoder and then globally pooled across all spatial and views dimensions to produce a scene-level descriptor $\boldsymbol{\theta}$:

$$\boldsymbol{\theta} = F_M(\{\mathbf{g}_v\}, \{\mathbf{e}_v\}, \mathcal{G}) \in \mathbb{R}^{d_\theta}, \quad (11)$$

where d_θ encodes all weights and biases of the two-layer tonemapper g_θ .

3.5 Training Strategies

Training Objective. InstantHDR is trained end-to-end without any 3D or HDR supervision, using only multi-view LDR images with known exposure times. The geometry encoder and its decoder heads remain frozen; only the appearance branch, the Gaussian head, and the Meta Net are optimized.

During training, the target views are the same with the context views, i.e., the model predicts HDR Gaussians $\mathcal{G}$ and tonemapper parameters $\boldsymbol{\theta}$ from $\{I_v, \ell_v\}_{v=1}^V$ and is supervised by rendering back to the same views and exposures via Eqs. (8)–(9). At test time, the target view and exposure can differ from the context set, enabling novel-view synthesis at arbitrary exposures. The total loss is:

$$\mathcal{L} = \mathcal{L}_{\text{RGB}} + \lambda_g \mathcal{L}_g. \quad (12)$$

The photometric loss compares rendered and ground-truth LDR images:

$$\mathcal{L}_{\text{RGB}} = \frac{1}{V} \sum_{v=1}^V \left[\text{MSE}(I_v, \mathbf{L}_v(\ell_v)) + \lambda_{\text{perc}} \cdot \mathcal{L}_{\text{perc}}(I_v, \mathbf{L}_v(\ell_v)) \right]. \quad (13)$$

The geometry consistency loss enforces alignment between the depth maps D_v from the frozen DPT head and the rendered depth maps $\hat{D}_v$ from the predicted Gaussians. Since D_v can be unreliable in challenging regions (e.g., sky or reflective surfaces), we utilize the jointly learned confidence map C_v^D and apply supervision only to the top- $N\%$ most confident pixels:

$$\mathcal{L}_g = \frac{1}{V} \sum_{v=1}^V (D_v[M_v] - \hat{D}_v[M_v])^2, \quad (14)$$

where M_v is a binary mask corresponding to the top- N quantile of C_v^D ; we set $N=30$ in all experiments. Our model learns HDR representations implicitly by correctly reconstructing LDR images across different exposure times.

Table 1: Comparison of existing HDR datasets for novel view synthesis. “Pano.” denotes panoramic 360° images. “Multi-Exp.” indicates multi-exposure LDR images, with the number of exposures in parentheses.

Dataset	Source	#Scenes	Type	HDR	GT	Multi-Exp.	Depth/Normal
HDR-NeRF [17]	Blender + Real	8 + 4	Object/Indoor	✓	✓	(5)	×
HDR-Plenoxels [25]	Blender + Real	5 + 4	Indoor	×	✓	(3)	×
Pano-NeRF [50]	Blender + Replica + Real	5 + 8 + 3	Indoor (Pano.)	✓	✓	(9)	✓
PanDORA [7]	Real capture (360°)	14	Indoor (Pano.)	✓	✓	(2)	×
HDR-Pretrain (Ours)	Blender	168	Indoor	✓	✓	(5)	✓

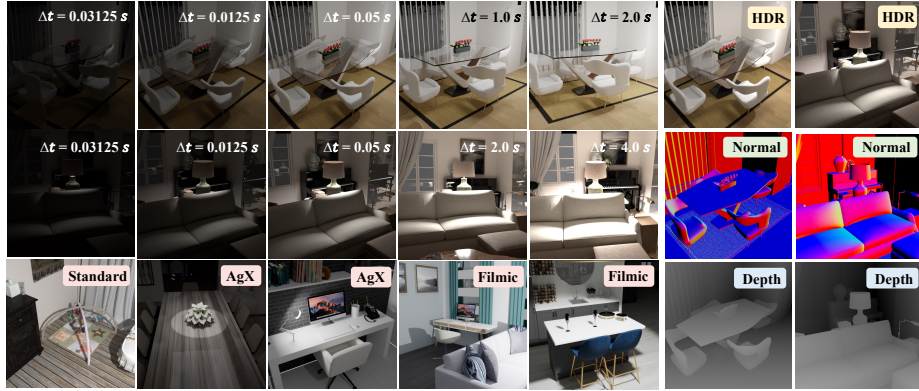


Fig. 4: Examples from our HDR-Pretrain dataset. Each scene includes multi-view, multi-exposure LDR images at varying Δt , 32-bit HDR ground truth, depth and normal maps, rendered under diverse tone-mapping operators (Standard, AgX, Filmic).

Post-Optimization. We can refine the predicted Gaussians and camera parameters via lightweight post-optimization. After pruning low-opacity Gaussians ($\sigma < 0.01$), we minimize a combination of MSE and SSIM losses between rendered and input images for 1K iterations, back-propagating gradients through all Gaussian attributes, and tonemapper parameters. The learning rates are: 1.6e-4 (position), 5e-3 (scale), 1e-3 (rotation), 5e-2 (opacity), and 2.5e-3 (color).

4 Experiments

Pretraining Dataset. HDR datasets with multi-view LDR images remain extremely scarce. As shown in Tab. 1, existing benchmarks contain only a handful of scenes, far from sufficient for large-scale pretraining. To bridge this gap, we construct **HDR-Pretrain**, in Fig. 4, a large-scale synthetic dataset of 168 photorealistic indoor scenes rendered in Blender. The 3D assets are sourced from HSSD [30], an open-source collection of realistic interiors originally built for embodied AI research. Following HDR-NeRF [17], we sample viewpoints on a 5×7 grid with $2.5^\circ / 5^\circ$ angular steps per scene and render 32-bit HDR images via Cycles path tracing at 448×448 resolution. Each view is paired with 5 exposure-bracketed LDR images under a randomly chosen tone-mapping operator, as well

as depth and normal maps. We randomly apply one of three tone-mapping operators (AgX, Filmic, Standard) per scene to increase data diversity.

Evaluation Dataset. We mainly evaluate on the HDR-NeRF benchmark [17], which contains 8 synthetic and 4 real indoor scenes, each captured from 35 view-points at 5 exposure levels $\{t_1, t_2, t_3, t_4, t_5\}$. Following the standard protocol [17], 18 views with one exposure drawn from $\{t_1, t_3, t_5\}$ are used as input, while the remaining views are held out for evaluation. For clarity, we report the average LDR metrics across all 5 exposures rather than separating observed (LDR-OE) and novel (LDR-NE) exposures. We further test under sparse-view settings with only 4 or 8 input views.

Implementation Details. We build upon AnySplat [21] with its backbone frozen, using a voxel size of $\epsilon=0.002$ in normalized scene coordinates. Training uses AdamW with cosine scheduling, peak learning rate 2×10^{-4} , 1K warmup, and runs for 30K iterations in bf16 precision on 8 NVIDIA A6000 GPUs for ~ 2 days. Input resolution is 448×448 with random cropping and flipping augmentation. We set $\lambda_{\text{perc}}=0.05$ and $\lambda_g=0.1$, sampling 2 $\sim$ 10 context views per iteration. Due to the domain gap between real and our synthetic HDR-Pretrain scenes, we finetune on HDR-Plenoxels 4 real scenes before evaluating on HDR-NeRF real scenes, ensuring all testing scenes are unseen. Following pose-free methods [88], we perform test-time pose alignment for evaluation and post-optimization.

Evaluation Metrics. All images are resized to 448×448 for fair comparison. We report PSNR, SSIM, LPIPS [95], and reconstruction time for quantitative and efficiency comparison. Following HDR-NeRF [17], HDR results are evaluated in the tone-mapped domain using the μ -law [26], with 99th-percentile normalization to suppress extreme HDR values.

4.1 Quantitative Results

LDR Comparisons on Zero-shot Inference. We first evaluate zero-shot results on HDR-NeRF [17] scenes. AnySplat [21] ignores exposure inconsistency, leading to severely degraded results. As shown in Tab. 2, InstantHDR consistently outperforms AnySplat by large margins (*e.g.*, +5.65 dB on real scenes, +8.07 dB on synthetic scenes with 8 views) with negligible increase in reconstruction time, showing strong generalization to exposure-inconsistent inputs.

LDR Comparisons on Optimization. We further compare against state-of-the-art optimization-based methods, HDR-GS [3] and GaussHDR [41]. In this setting, we post-optimize the Gaussians and tone-mapping parameters produced by InstantHDR and AnySplat for 1K iterations, denoted as InstantHDR_1K and AnySplat_1K. Under the challenging 4-view sparse setting on real scenes, InstantHDR_1K achieves 22.16 dB PSNR and 0.762 SSIM, surpassing GaussHDR by +2.90 dB and +0.071 in SSIM, demonstrating that the diverse geometric priors from feed-forward foundation models greatly benefit sparse-view reconstruction. Under the denser 18-view setting, our method remains competitive on real scenes while showing a modest gap on synthetic scenes. Notably, InstantHDR_1K requires only ~ 30 – 40 seconds per scene, roughly **20 $\times$** faster than

Table 2: Quantitative LDR comparison on HDR-NeRF [17] real and synthetic scenes with varying numbers of input views. We compare our method, as well as its variants with 1K iters of post-optimization (Ours_1K), against AnySplat [21] HDR-GS [3] and GaussHDR [41]. Time denotes per-scene reconstruction time in seconds.

Mode	Method	4 Views				8 Views				18 Views			
		PSNR↑	SSIM↑	LPIPS↓	Time(s)↓	PSNR↑	SSIM↑	LPIPS↓	Time(s)↓	PSNR↑	SSIM↑	LPIPS↓	Time(s)↓
<i>HDR-NeRF Real Dataset [17]</i>													
Zero-shot	AnySplat [21]	12.10	0.517	0.497	0.973	13.30	0.569	0.436	1.180	13.91	0.600	0.403	2.139
	InstantHDR (Ours)	18.44	0.721	0.269	1.033	18.95	0.724	0.269	1.582	19.48	0.745	0.257	2.512
Optimization	HDR-GS [3]	15.40	0.622	0.334	872	23.02	0.791	0.121	736	27.42	0.893	0.047	815
	GaussHDR [41]	19.26	0.691	0.270	1833	24.96	0.854	0.068	1816	29.36	0.929	0.024	1891
	AnySplat_1K [21]	11.84	0.486	0.468	30	13.63	0.580	0.372	33	14.86	0.689	0.266	40
	InstantHDR_1K (Ours)	22.16	0.762	0.259	32	25.32	0.852	0.160	40	29.19	0.931	0.086	39
<i>HDR-NeRF Syn Dataset [17]</i>													
Zero-shot	AnySplat [21]	13.98	0.525	0.400	-	14.51	0.573	0.378	-	15.27	0.534	0.327	-
	InstantHDR (Ours)	21.76	0.728	0.172	-	22.58	0.785	0.138	-	22.59	0.830	0.115	-
Optimization	HDR-GS [3]	24.26	0.711	0.210	-	30.60	0.867	0.086	-	29.93	0.917	0.061	-
	GaussHDR [41]	21.62	0.646	0.224	-	34.49	0.924	0.026	-	38.63	0.969	0.009	-
	AnySplat_1K [21]	14.01	0.526	0.347	-	15.42	0.656	0.248	-	16.39	0.624	0.221	-
	InstantHDR_1K (Ours)	27.63	0.825	0.137	-	32.75	0.922	0.061	-	35.99	0.965	0.037	-

Table 3: (a) Quantitative HDR comparison on HDR-NeRF [17] synthetic scenes with 8 input views, evaluated in the μ -law tone-mapped domain. (b) Ablation study of InstantHDR. Zero-shot results on HDR-NeRF [17] real scenes with 8 input views.

(a) HDR Comparison on HDR-NeRF [17] syn scenes.						(b) Ablation on HDR-NeRF [17] real scenes			
Mode	Method	PSNR↑	SSIM↑	LPIPS↓	Time(s)↓	Method	PSNR↑	SSIM↑	LPIPS↓
Zero-shot	AnySplat	8.93	0.595	0.416	1.180	w/o Meta Net	16.32	0.699	0.289
	InstantHDR (Ours)	15.29	0.772	0.140	1.582	w/o Exposure Norm	13.72	0.693	0.278
Optimization	HDR-GS	27.69	0.871	0.090	736	w/o Cross-view Attn	17.63	0.702	0.277
	GaussHDR	31.62	0.887	0.037	1816	w/o Upsampling	19.20	0.718	0.386
	AnySplat_1K	9.52	0.678	0.268	33	Ours	18.95	0.724	0.269
	InstantHDR_1K (Ours)	27.55	0.899	0.076	40				

HDR-GS and 50× faster than GaussHDR, attributing to the good feed-forward initialization and the ability to skip the costly iterative densification.

HDR Comparisons. Tab. 3 (a) reports HDR results on HDR-NeRF synthetic scenes. Although never supervised with HDR ground truth, InstantHDR implicitly recovers HDR from multi-exposure LDR inputs, outperforming AnySplat by +6.36 dB in PSNR in zero-shot mode. With 1K steps of post-optimization, InstantHDR_1K reaches 27.55 dB PSNR and 0.899 SSIM, achieving comparable performance to HDR-GS (27.69 dB) while surpassing GaussHDR in SSIM. The remaining gap with GaussHDR in PSNR is likely due to its dedicated 3D–2D dual-branch tone-mapping design, whereas both our method and HDR-GS adopt a simpler single tone mapping branch. We believe that incorporating more advanced tone-mapping modules is a promising avenue for future works.

4.2 Qualitative Results

LDR Novel View Rendering. In Fig. 5, AnySplat struggles in exposure variations, while GaussHDR is time-consuming. Our InstantHDR generalizes well to both synthetic and real-world scenes, renders clean, exposure-controllable LDR views in seconds, and achieves competitive quality after 1K post-optimization.



(a) LDR visual comparisons on HDR-NeRF [17] synthetic scenes.



(b) LDR visual comparisons on HDR-NeRF [17] real scenes.

Fig. 5: LDR visual comparisons. Feed-forward methods [21] fail on multi-exposure inputs, while optimization-based methods [41] require $\sim 2K$ seconds per scene. Our InstantHDR achieves competitive quality in under 40s. Yellow/blue tags denote reconstruction time/PSNR.

HDR Novel View Rendering. As shown in Fig. 6, zero-shot HDR outputs from feed-forward models (AnySplat & InstantHDR) appear overly bright, as extreme radiance values are hard to predict accurately in a single-forward pass, which inflates the average brightness after normalization. We see this as an open challenge for feed-forward HDR models. After 1K post-optimization, this issue is largely alleviated, with our InstantHDR producing results similar to GaussHDR.



Fig. 6: HDR visual comparisons. After 1K post-optimization, our InstantHDR produces HDR results comparable to time-consuming GaussHDR.



Fig. 7: Ablation visualization. Attn.=cross-view attention, Exp.=exposure normalization, Up.=upsampling. Removing Attn. causes wall ghosting; removing Exp. shifts overall brightness; removing Up. produces blurry results.

4.3 Ablation Study

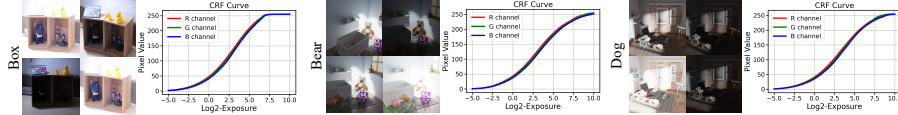
Baseline. We show that naïve extensions of existing pipelines fail. In Tab. 4 (a), we add exposure-aware tokens to AnySplat [21] but it performs limited: *its encoder is geometry-dominated, not color; color is only introduced via skip connection, which lacks cross-view interaction to resolve exposure inconsistency.* Our key insight is that the frozen geometry encoder already captures reliable cross-view correspondences under large exposure gaps, and its attention maps can be reused for HDR appearance fusion, this has not been explored before.

Component ablation. Tab. 3 (b) and Fig. 7 confirm each component is critical. Removing the MetaNet destabilizes training, as the model cannot adapt to varying camera response functions. Eliminating exposure normalization causes the largest degradation, since inconsistent brightness across views disrupts feature fusion. Disabling cross-view attention introduces ghosting on smooth surfaces such as walls, confirming the value of reusing the geometry encoder’s correspondences. Removing the upsampling module preserves coarse structure but loses fine details, yielding blurry outputs.

CRF curves. To further illustrate the MetaNet ablated above, we show that it automatically learns monotonic, scene-specific CRF curves through large-scale

Table 4: (a) Comparison against an exposure-aware AnySplat (HDR) baseline. (b) Effect of one-time real-scene finetuning.

(a) Compared with AnySplat (HDR) baseline.							(b) One-time real-scene finetuning.						
HDR-NeRF Real	4 Views			18 Views			HDR-NeRF Real	4 Views			18 Views		
	PSNR $\uparrow$	SSIM $\uparrow$	LPIPS $\downarrow$	PSNR $\uparrow$	SSIM $\uparrow$	LPIPS $\downarrow$		PSNR $\uparrow$	SSIM $\uparrow$	LPIPS $\downarrow$	PSNR $\uparrow$	SSIM $\uparrow$	LPIPS $\downarrow$
AnySplat (HDR)	14.59	0.571	0.419	15.65	0.629	0.377	w/o finetuning	17.27	0.704	0.298	18.08	0.715	0.271
InstantHDR	18.44	0.721	0.269	19.48	0.745	0.257	w finetuning	18.44	0.721	0.269	19.48	0.745	0.257

**Fig. 8: CRF visualization.** Learned CRF curves for three scenes, all monotonic.

pretraining. As shown in Fig. 8, the Box scene saturates earlier while others span a wider dynamic range, all strictly monotonic.

Real-scene finetuning. Since the pretraining dataset is synthetic, we perform a *one-time finetuning* on 4 HDR-Plenoxels real scenes to bridge the synthetic-to-real gap—not per-scene optimization and no test-data leaking—and apply this model to all real scenes without further tuning. As shown in Tab. 4 (b), zero-shot (no finetuning) results still generalize well, while finetuning adds ~ 1.28 dB.

5 Conclusion

We present InstantHDR, the first feed-forward framework for HDR novel view synthesis from uncalibrated multi-exposure LDR images. By leveraging geometry-guided cross-view attention for exposure-robust appearance fusion and a meta-network for scene-adaptive tone mapping, InstantHDR initializes HDR scenes in few seconds. We also introduce HDR-Pretrain, a 168-scene synthetic dataset to address the data scarcity for feed-forward HDR pretraining. Experiments show that, with post-optimization, InstantHDR reaches quality comparable to optimization-based methods at orders-of-magnitude-faster speed. We hope this InstantHDR can inspire more ideas on real-time 3D HDR reconstruction.

Limitation. Our InstantHDR focuses on HDR appearance modeling and relies on feed-forward backbones for pose estimation. On noise-free dense synthetic scenes, our feed-forward pose estimator is less accurate than COLMAP, which limits reconstruction quality—a limitation shared by all feed-forward methods.

References

1. Bolduc, C., Hold-Geoffroy, Y., Shu, Z., Lalonde, J.F.: Gaslight: Gaussian splats for spatially-varying lighting in hdr. In: ICCV (2025)
2. Cai, Y., Liang, Y., Wang, J., Wang, A., Zhang, Y., Yang, X., Zhou, Z., Yuille, A.: Radiative gaussian splatting for efficient x-ray novel view synthesis. In: ECCV (2024)

3. Cai, Y., Xiao, Z., Liang, Y., Qin, M., Zhang, Y., Yang, X., Liu, Y., Yuille, A.L.: Hdr-gs: Efficient high dynamic range novel view synthesis at 1000x speed via gaussian splatting. *NeurIPS* **37**, 68453–68471 (2024)
4. Chen, R., Zheng, B., Zhang, H., Chen, Q., Yan, C., Slabaugh, G., Yuan, S.: Improving dynamic hdr imaging with fusion transformer. In: *AAAI* (2023)
5. Chen, Z., Yang, J., Yang, H.: Pref3r: Pose-free feed-forward 3d gaussian splatting from variable-length image sequence. *arXiv preprint arXiv:2411.16877* (2024)
6. Cui, Z., Chu, X., Harada, T.: Luminance-gs: Adapting 3d gaussian splatting to challenging lighting conditions with view-adaptive curve adjustment. In: *CVPR*. pp. 26472–26482 (2025)
7. Dastjerdi, M.R.K., Tanguay-Gaudreau, D., Fortier-Chouinard, F., Hold-Geoffroy, Y., Demers, C., Kalantari, N., Lalonde, J.F.: Pandora: Casual hdr radiance acquisition for indoor scenes. *arXiv preprint arXiv:2407.06150* (2024)
8. Debevec, P.E., Malik, J.: Recovering high dynamic range radiance maps from photographs. In: *SIGGRAPH* (1997)
9. Eilertsen, G., Kronander, J., Denes, G., Mantiuk, R.K., Unger, J.: Hdr image reconstruction from a single exposure using deep cnns. *ACM TOG* (2017)
10. Fei, B., Lyu, Z., Pan, L., Zhang, J., Yang, W., Luo, T., Zhang, B., Dai, B.: Generative diffusion prior for unified image restoration and enhancement. In: *CVPR* (2023)
11. Gong, S., Zhao, L., Li, W., Xie, H., Zhang, Y., Zhao, S., Liu, P.: Casual3dhdr: Deblurring high dynamic range 3d gaussian splatting from casually captured videos. *ACMMM* (2025)
12. Grosch, T., et al.: Fast and robust high dynamic range image generation with camera and object movement. *Vision, Modeling and Visualization, RWTH Aachen* (2006)
13. Hong, S., Jung, J., Shin, H., Han, J., Yang, J., Luo, C., Kim, S.: Pf3plat: Pose-free feed-forward 3d gaussian splatting. *arXiv preprint arXiv:2410.22128* (2024)
14. Hu, S., Hu, T., Liu, Z.: Gauhuman: Articulated gaussian splatting from monocular human videos. In: *CVPR*. pp. 20418–20431 (2024)
15. Hu, T., Sarkar, K., Liu, L., Zwicker, M., Theobalt, C.: Egorenderer: Rendering human avatars from egocentric camera images. In: *ICCV*. pp. 14528–14538 (2021)
16. Huang, S., Gojcic, Z., Wang, Z., Williams, F., Kasten, Y., Fidler, S., Schindler, K., Litany, O.: Neural lidar fields for novel view synthesis. In: *ICCV*. pp. 18236–18246 (2023)
17. Huang, X., Zhang, Q., Feng, Y., Li, H., Wang, X., Wang, Q.: Hdr-nerf: High dynamic range neural radiance fields. In: *CVPR*. pp. 18398–18408 (2022)
18. Jacobs, K., Loscos, C., Ward, G.: Automatic high-dynamic range image generation for dynamic scenes. *IEEE Computer Graphics and Applications* (2008)
19. Jiang, H., Guan, B., Liu, Z., Liu, X., Yu, J., Liu, Z., Han, S., Liu, S.: Learning to see in the extremely dark. In: *ICCV* (2025)
20. Jiang, H., Jiang, Z., Zhao, Y., Huang, Q.: Leap: Liberate sparse-view 3d modeling from camera poses. *arXiv preprint arXiv:2310.01410* (2023)
21. Jiang, L., Mao, Y., Xu, L., Lu, T., Ren, K., Jin, Y., Xu, X., Yu, M., Pang, J., Zhao, F., et al.: Anysplat: Feed-forward 3d gaussian splatting from unconstrained views. *ACM TOG* **44**(6), 1–16 (2025)
22. Jiang, Y., Tu, J., Liu, Y., Gao, X., Long, X., Wang, W., Ma, Y.: Gaussianshader: 3d gaussian splatting with shading functions for reflective surfaces. *arXiv preprint arXiv:2311.17977* (2023)
23. Jin, H., Li, Y., Luan, F., Xiangli, Y., Bi, S., Zhang, K., Xu, Z., Sun, J., Snavely, N.: Neural gaffer: Relighting any object via diffusion. In: *NeurIPS* (2024)

24. Jin, X., Jiao, P., Duan, Z.P., Yang, X., Li, C., Guo, C.L., Ren, B.: Lighting every darkness with 3dgs: Fast training and real-time rendering for hdr view synthesis. *NeurIPS* **37**, 80191–80219 (2024)
25. Jun-Seong, K., Yu-Ji, K., Ye-Bin, M., Oh, T.H.: Hdr-plenoxels: Self-calibrating high dynamic range radiance fields. In: *ECCV*. pp. 384–401. Springer (2022)
26. Kalantari, N.K., Ramamoorthi, R., et al.: Deep high dynamic range imaging of dynamic scenes. *ACM ToG* (2017)
27. Keetha, N., Karhade, J., Jatavallabhula, K.M., Yang, G., Scherer, S., Ramanan, D., Luiten, J.: Splatam: Splat, track & map 3d gaussians for dense rgb-d slam. *arXiv preprint arXiv:2312.02126* (2023)
28. Kerbl, B., Kopanas, G., Leimkühler, T., Drettakis, G.: 3d gaussian splatting for real-time radiance field rendering. *ACM Transactions on Graphics* (2023)
29. Khan, Z., Khanna, M., Raman, S.: Fhdr: Hdr image reconstruction from a single ldr image using feedback network. In: *IEEE Global Conference on Signal and Information Processing* (2019)
30. Khanna, M., Mao, Y., Jiang, H., Hares, S., Shacklett, B., Batra, D., Clegg, A., Undersander, E., Chang, A.X., Savva, M.: Habitat synthetic scenes dataset (hssd-200): An analysis of 3d scene scale and realism tradeoffs for objectgoal navigation. In: *CVPR*. pp. 16384–16393 (2024)
31. Kim, J., Lee, S., Kang, S.J.: End-to-end differentiable learning to hdr image synthesis for multi-exposure images. In: *AAAI* (2021)
32. Kobayashi, S., Matsumoto, E., Sitzmann, V.: Decomposing nerf for editing via feature field distillation. *NeurIPS* **35**, 23311–23330 (2022)
33. Kocabas, M., Chang, J.H.R., Gabriel, J., Tuzel, O., Ranjan, A.: Hugs: Human gaussian splats. *arXiv preprint arXiv:2311.17910* (2023)
34. Lee, S., Park, E., Canelo, A., Park, H., Kim, Y., Chun, H., Jin, X., Li, C., Guo, C.L., Timofte, R., et al.: Ntire 2025 challenge on efficient burst hdr and restoration: Datasets, methods, and results. In: *CVPRW*. pp. 1002–1017 (2025)
35. Leroy, V., Cabon, Y., Revaud, J.: Grounding image matching in 3d with mast3r. In: *ECCV*. pp. 71–91. Springer (2024)
36. Li, R., Wang, Y., Chen, S., Zhang, F., Gu, J., Xue, T.: Dualdn: Dual-domain denoising via differentiable isp. In: *ECCV* (2024)
37. Li, Y., Wang, H., Xu, K., Hancke, G.P., Lau, R.W.: Sehdr: Single-exposure hdr novel view synthesis via 3d gaussian bracketing. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision*. pp. 26045–26054 (2025)
38. Li, Z., Wang, Y., Kot, A., Wen, B.: From chaos to clarity: 3dgs in the dark. In: *NeurIPS*. vol. 37, pp. 94971–94992 (2024)
39. Liang, Y., Yang, X., Lin, J., Li, H., Xu, X., Chen, Y.: Luciddreamer: Towards high-fidelity text-to-3d generation via interval score matching. *arXiv preprint arXiv:2311.11284* (2023)
40. Liang, Z., Zhang, Q., Feng, Y., Shan, Y., Jia, K.: Gs-ir: 3d gaussian splatting for inverse rendering. *arXiv preprint arXiv:2311.16473* (2023)
41. Liu, J., Kong, L., Li, B., Xu, D.: Gausshdr: High dynamic range gaussian splatting via learning unified 3d and 2d local tone mapping. In: *CVPR*. pp. 5991–6000 (2025)
42. Liu, J., Kong, L., Zhou, M., Chen, J., Xu, D.: Mono4dgs-hdr: High dynamic range 4d gaussian splatting from alternating-exposure monocular videos. In: *ICLR* (2026)
43. Liu, L., Habermann, M., Rudnev, V., Sarkar, K., Gu, J., Theobalt, C.: Neural actor: Neural free-view synthesis of human actors with pose control. *ACM TOG* **40**(6), 1–16 (2021)

44. Liu, S., Zhang, X., Sun, L., Liang, Z., Zeng, H., Zhang, L.: Joint hdr denoising and fusion: A real-world mobile hdr image dataset. In: CVPR (2023)
45. Liu, S., Zhang, X., Zhang, Z., Zhang, R., Zhu, J.Y., Russell, B.: Editing conditional radiance fields. In: ICCV. pp. 5773–5783 (2021)
46. Liu, X., Zhan, X., Tang, J., Shan, Y., Zeng, G., Lin, D., Liu, X., Liu, Z.: Humangaussian: Text-driven 3d human generation with gaussian splatting. arXiv preprint arXiv:2311.17061 (2023)
47. Liu, Y., Dong, S., Wang, S., Yin, Y., Yang, Y., Fan, Q., Chen, B.: Slam3r: Real-time dense scene reconstruction from monocular rgb videos. In: CVPR. pp. 16651–16662 (2025)
48. Liu, Z., Lin, W., Li, X., Rao, Q., Jiang, T., Han, M., Fan, H., Sun, J., Liu, S.: Adnet: Attention-guided deformable convolutional network for high dynamic range imaging. In: CVPRW (2021)
49. Liu, Z., Wang, Y., Zeng, B., Liu, S.: Ghost-free high dynamic range imaging with context-aware transformer. In: ECCV (2022)
50. Lu, Z., Zheng, Q., Shi, B., Jiang, X.: Pano-nerf: Synthesizing high dynamic range novel views with geometry from sparse low dynamic range panoramic images. In: AAAI. pp. 3927–3935 (2024)
51. Luiten, J., Kopanas, G., Leibe, B., Ramanan, D.: Dynamic 3d gaussians: Tracking by persistent dynamic view synthesis. arXiv preprint arXiv:2308.09713 (2023)
52. Matsuki, H., Murai, R., Kelly, P.H., Davison, A.J.: Gaussian splatting slam. arXiv preprint arXiv:2312.06741 (2023)
53. Mertens, T., Kautz, J., Van Reeth, F.: Exposure fusion. In: Pacific Conference on Computer Graphics and Applications (2007)
54. Mildenhall, B., Hedman, P., Martin-Brualla, R., Srinivasan, P.P., Barron, J.T.: Nerf in the dark: High dynamic range view synthesis from noisy raw images. In: CVPR (2022)
55. Mildenhall, B., Srinivasan, P.P., Tancik, M., Barron, J.T., Ramamoorthi, R., Ng, R.: Nerf: Representing scenes as neural radiance fields for view synthesis. *Communications of the ACM* **65**(1), 99–106 (2021)
56. Murai, R., Dexheimer, E., Davison, A.J.: Mast3r-slam: Real-time dense slam with 3d reconstruction priors. In: CVPR. pp. 16695–16705 (2025)
57. Nazarczuk, M., Catley-Chandar, S., Leonardis, A., Pellitero, E.P.: Self-supervised hdr imaging from motion and exposure cues. arXiv preprint arXiv:2203.12311 (2022)
58. Oquab, M., Darcet, T., Moutakanni, T., Vo, H., Szafraniec, M., Khalidov, V., Fernandez, P., Haziza, D., Massa, F., El-Nouby, A., et al.: Dinov2: Learning robust visual features without supervision. arXiv preprint arXiv:2304.07193 (2023)
59. Perez, E., Strub, F., De Vries, H., Dumoulin, V., Courville, A.: Film: Visual reasoning with a general conditioning layer. In: AAAI (2018)
60. Prabhakar, K.R., Senthil, G., Agrawal, S., Babu, R.V., Gorthi, R.K.S.S.: Labeled from unlabeled: Exploiting unlabeled data for few-shot deep hdr deghosting. In: CVPR (2021)
61. Ranftl, R., Bochkovskiy, A., Koltun, V.: Vision transformers for dense prediction. In: ICCV. pp. 12179–12188 (2021)
62. Shu, Y., Shen, L., Hu, X., Li, M., Zhou, Z.: Towards real-world hdr video reconstruction: A large-scale benchmark dataset and a two-stage alignment network. In: CVPR (2024)
63. Singh, S., Garg, A., Mitra, K.: Hdrsplat: Gaussian splatting for high dynamic range 3d scene reconstruction from raw images. *BMVC* (2024)

64. Smart, B., Zheng, C., Laina, I., Prisacariu, V.A.: Splatt3r: Zero-shot gaussian splatting from uncalibrated image pairs. *arXiv preprint arXiv:2408.13912* (2024)
65. Song, J.W., Park, Y.I., Kong, K., Kwak, J., Kang, S.J.: Selective transhdr: Transformer-based selective hdr imaging using ghost region mask. In: *ECCV* (2022)
66. Sun, J., Wang, X., Shi, Y., Wang, L., Wang, J., Liu, Y.: Ide-3d: Interactive disentangled editing for high-resolution 3d-aware portrait synthesis. *ACM TOG* **41**(6), 1–10 (2022)
67. Tancik, M., Casser, V., Yan, X., Pradhan, S., Mildenhall, B., Srinivasan, P.P., Barron, J.T., Kretzschmar, H.: Block-nerf: Scalable large scene neural view synthesis. In: *CVPR*. pp. 8248–8258 (2022)
68. Tang, J., Ren, J., Zhou, H., Liu, Z., Zeng, G.: Dreamgaussian: Generative gaussian splatting for efficient 3d content creation. *arXiv preprint arXiv:2309.16653* (2023)
69. Tang, Z., Fan, Y., Wang, D., Xu, H., Ranjan, R., Schwing, A., Yan, Z.: Mvdust3r+: Single-stage scene reconstruction from sparse views in 2 seconds. In: *CVPR*. pp. 5283–5293 (2025)
70. Tursun, O.T., Akyüz, A.O., Erdem, A., Erdem, E.: The state of the art in hdr deghosting: A survey and evaluation. In: *Computer Graphics Forum* (2015)
71. Wang, G., Chen, Z., Loy, C.C., Liu, Z.: Sparsenerf: Distilling depth ranking for few-shot novel view synthesis. In: *ICCV*. pp. 9065–9076 (2023)
72. Wang, H., Agapito, L.: 3d reconstruction with spatial memory. *arXiv preprint arXiv:2408.16061* (2024)
73. Wang, J., Chen, M., Karaev, N., Vedaldi, A., Rupperecht, C., Novotny, D.: Vggt: Visual geometry grounded transformer. In: *CVPR*. pp. 5294–5306 (2025)
74. Wang, P., Tan, H., Bi, S., Xu, Y., Luan, F., Sunkavalli, K., Wang, W., Xu, Z., Zhang, K.: Pf-lrm: Pose-free large reconstruction model for joint pose and shape prediction. *arXiv preprint arXiv:2311.12024* (2023)
75. Wang, Q., Zhang, Y., Holynski, A., Efros, A.A., Kanazawa, A.: Continuous 3d perception model with persistent state. *arXiv preprint arXiv:2501.12387* (2025)
76. Wang, S., Leroy, V., Cabon, Y., Chidlovskii, B., Revaud, J.: Dust3r: Geometric 3d vision made easy. In: *CVPR*. pp. 20697–20709 (2024)
77. Ward, G., Reinhard, E., Debevec, P.: High dynamic range imaging & image-based lighting. In: *SIGGRAPH* (2008)
78. Wu, G., Yi, T., Fang, J., Liu, W., Wang, X.: Fast high dynamic range radiance fields for dynamic scenes. In: *2024 International Conference on 3D Vision (3DV)*. pp. 862–872. *IEEE* (2024)
79. Wu, G., Yi, T., Fang, J., Xie, L., Zhang, X., Wei, W., Liu, W., Tian, Q., Wang, X.: 4d gaussian splatting for real-time dynamic scene rendering. *arXiv preprint arXiv:2310.08528* (2023)
80. Xie, T., Zong, Z., Qiu, Y., Li, X., Feng, Y., Yang, Y., Jiang, C.: Physgaussian: Physics-integrated 3d gaussians for generative dynamics. *arXiv preprint arXiv:2311.12198* (2023)
81. Xu, G., Wang, Y., Gu, J., Xue, T., Yang, X.: Hdrflow: Real-time hdr video reconstruction with large motions. In: *CVPR* (2024)
82. Yan, C., Qu, D., Wang, D., Xu, D., Wang, Z., Zhao, B., Li, X.: Gs-slam: Dense visual slam with 3d gaussian splatting. *arXiv preprint arXiv:2311.11700* (2023)
83. Yan, Q., Zhang, S., Chen, W., Tang, H., Zhu, Y., Sun, J., Van Gool, L., Zhang, Y.: Smae: Few-shot learning for hdr deghosting with saturation-aware masked autoencoders. In: *CVPR* (2023)
84. Yan, Q., Zhu, Y., Zhang, Y.: Robust artifact-free high dynamic range imaging of dynamic scenes. *Multimedia Tools and Applications* (2019)

85. Yang, J., Sax, A., Liang, K.J., Henaff, M., Tang, H., Cao, A., Chai, J., Meier, F., Feiszli, M.: Fast3r: Towards 3d reconstruction of 1000+ images in one forward pass. arXiv preprint arXiv:2501.13928 (2025)
86. Yang, Z., Yang, H., Pan, Z., Zhu, X., Zhang, L.: Real-time photorealistic dynamic scene representation and rendering with 4d gaussian splatting. arXiv preprint arXiv:2310.10642 (2023)
87. Yang, Z., Chai, Y., Anguelov, D., Zhou, Y., Sun, P., Erhan, D., Rafferty, S., Kretschmar, H.: Surfelgan: Synthesizing realistic sensor data for autonomous driving. In: CVPR. pp. 11118–11127 (2020)
88. Ye, B., Liu, S., Xu, H., Li, X., Pollefeys, M., Yang, M.H., Peng, S.: No pose, no problem: Surprisingly simple 3d gaussian splats from sparse unposed images. arXiv preprint arXiv:2410.24207 (2024)
89. Yi, T., Fang, J., Wu, G., Xie, L., Zhang, X., Liu, W., Tian, Q., Wang, X.: Gaussianreamer: Fast generation from text to 3d gaussian splatting with point cloud priors. arXiv preprint arXiv:2310.08529 (2023)
90. Yu, Y., Wang, H., Luo, T., Fan, H., Zhang, L.: Magic: Multi-modality guided image completion. In: ICLR (2024)
91. Yu, Y., Zeng, Z., Hua, H., Fu, J., Luo, J.: Promptfix: You prompt and we fix the photo. In: NeurIPS (2024)
92. Yuan, Y.J., Sun, Y.T., Lai, Y.K., Ma, Y., Jia, R., Gao, L.: Nerf-editing: geometry editing of neural radiance fields. In: CVPR. pp. 18353–18364 (2022)
93. Yugay, V., Li, Y., Gevers, T., Oswald, M.R.: Gaussian-slam: Photo-realistic dense slam with gaussian splatting. arXiv preprint arXiv:2312.10070 (2023)
94. Zha, R., Lin, T.J., Cai, Y., Cao, J., Zhang, Y., Li, H.: R²-gaussian: Rectifying radiative gaussian splatting for tomographic reconstruction. In: NeurIPS (2024)
95. Zhang, R., Isola, P., Efros, A.A., Shechtman, E., Wang, O.: The unreasonable effectiveness of deep features as a perceptual metric. In: CVPR. pp. 586–595 (2018)
96. Zhang, S., Wang, J., Xu, Y., Xue, N., Rupprecht, C., Zhou, X., Shen, Y., Wetstein, G.: Flare: Feed-forward geometry, appearance and camera estimation from uncalibrated sparse views. arXiv preprint arXiv:2502.12138 (2025)
97. Zhang, Z., Wang, H., Liu, S., Wang, X., Lei, L., Zuo, W.: Self-supervised high dynamic range imaging with multi-exposure images in dynamic scenes. In: ICLR (2024)
98. Zheng, J., Jang, Y., Papaioannou, A., Kampouris, C., Potamias, R.A., Papantoniou, F.P., Galanakis, E., Leonardis, A., Zafeiriou, S.: Ilsh: The imperial light-stage head dataset for human head view synthesis. In: ICCV. pp. 1112–1120 (2023)
99. Zheng, Z., Huang, H., Yu, T., Zhang, H., Guo, Y., Liu, Y.: Structured local radiance fields for human avatar modeling. In: CVPR. pp. 15893–15903 (2022)
100. Zhou, H., Dong, W., Chen, J.: Lita-gs: Illumination-agnostic novel view synthesis via reference-free 3d gaussian splatting and physical priors. In: CVPR. pp. 21580–21589 (2025)
101. Zou, Y., Yan, C., Fu, Y.: Rawhdr: High dynamic range image reconstruction from a single raw image. In: ICCV (2023)