

Relaxed Rigidity with Ray-based Grouping for Dynamic Gaussian Splatting

Junoh Lee^{1*}, Junmyeong Lee², Yeon-Ji Song³, Inhwan Bae⁴,
Jisu Shin⁵, Hae-Gon Jeon^{2†}, and Jin-Hwa Kim^{6,7†}

¹Department of EECS, GIST ²Department of Artificial Intelligence, Yonsei Univ.

³Department of CSE, SNU ⁴Department of EECS, DGIST

⁵Department of AI Convergence, GIST ⁶NAVER AI Lab ⁷AI Institute of SNU

{juno, jsshin98}@gm.gist.ac.kr, {june65, earboll}@yonsei.ac.kr,
yjsong@snu.ac.kr, inhwanbae@dgist.ac.kr, jinhwa.kim@navercorp.com

Abstract. The reconstruction of dynamic 3D scenes using 3D Gaussian Splatting has shown significant promise. A key challenge, however, remains in modeling realistic motion, as most methods fail to align the motion of Gaussians with real-world physical dynamics. This misalignment is particularly problematic for monocular video datasets, where failing to maintain coherent motion undermines local geometric structure, ultimately leading to degraded reconstruction quality. Consequently, many state-of-the-art approaches rely heavily on external priors, such as optical flow or 2D tracks, to enforce temporal coherence. In this work, we propose a novel method to explicitly preserve the local geometric structure of Gaussians across time in 4D scenes. Our core idea is to introduce a view-space ray grouping strategy that clusters Gaussians intersected by the same ray, considering only those whose α -blending weights exceed a threshold. We then apply constraints to these groups to maintain a consistent spatial distribution, effectively preserving their local geometry. This approach enforces a more locally coherent motion model by ensuring that local geometry remains stable over time, eliminating the reliance on external guidance. We demonstrate the efficacy of our method by integrating it into two distinct baseline models. Extensive experiments on challenging monocular datasets show that our approach significantly outperforms existing methods, achieving superior temporal consistency and reconstruction quality.

Keywords: Dynamic 3D Gaussian Splatting · Point Grouping · Motion Regularization

1 Introduction

Reconstructing dynamic scenes from videos is essential for modeling spatio-temporal structure in real-world scenes. The emergence of 3D Gaussian Splatting (3DGS) [18] has significantly advanced this task with its explicit representation and efficient training and rendering. Various dynamic extensions of 3DGS [26, 58,

* Work done during an internship at NAVER AI Lab.

† Corresponding authors

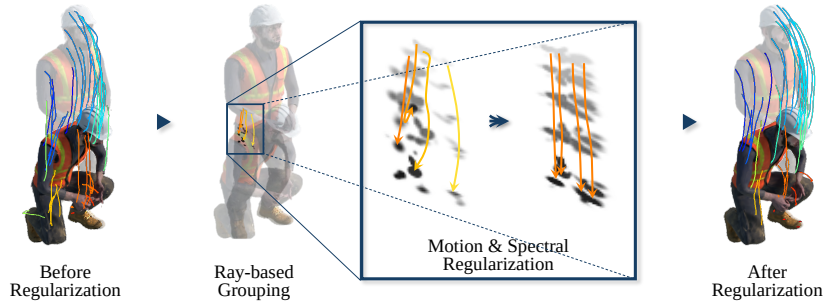


Fig. 1: Overview. Before regularization (first), individual Gaussians exhibit temporally inconsistent motion; each colored trajectory denotes the time-varying movement of a single Gaussian. Our method first performs ray-based grouping to cluster spatially adjacent Gaussians along view rays (second), with a zoomed-in example shown in the third panel. Motion and spectral regularization are then applied within each group to enforce coherent dynamics. After regularization (fourth), the Gaussian trajectories become temporally aligned and physically consistent.

66] have demonstrated impressive reconstruction, even under monocular camera settings. For instance, these approaches enable continuous spatio-temporal scene modeling by parameterizing time-varying Gaussian primitives.

However, accurately reflecting the physical motion of the real 4D world remains challenging, as the estimated motion is often under- or over-constrained and not directly governed by physical plausibility. To mitigate this issue, prior works [4, 72] employ motion priors derived from optical flow and depth or impose rigidity assumptions [12, 19, 28]. Nevertheless, these approaches exhibit fundamental limitations. External prior models [54, 61] are not inherently designed to guide the point dynamics of Gaussian primitives. Specifically, these methods project Gaussian motion onto the image plane and rely on optical or 3D flow as proxy guidance. Because such supervision is defined in 2D screen space rather than in the underlying 3D geometry, it provides only indirect and potentially inconsistent motion cues. As a result, errors and ambiguities in the proxy signals are propagated to the optimization process, leading to inaccurate learning of Gaussian motion [61, 72]. In contrast, rigid-motion-based models use K-Nearest Neighbors (KNN) to group adjacent points and enforce shared rigid transformations. This strategy overlooks the inherently non-rigid and interactive nature of real-world motion and fails to account for varying spatial scales of Gaussian primitives, which require adaptive grouping regions. Consequently, relying on external priors or strict rigidity assumptions is insufficient to ensure physically consistent 4D representations.

To enforce locally coherent motion constraints, we focus on two key aspects (Fig. 1). First, we group spatially adjacent Gaussians that exhibit coherent motion using a ray-based strategy rather than distance-based grouping. For each pixel, we consider only Gaussians intersected by the same view ray whose α -blending weights [18] exceed a threshold, excluding irrelevant primitives along the ray. This corresponds to the sorted and blended Gaussians at the pixel in 3DGS. By selecting only visible, high-contribution Gaussians, we obtain motion-

consistent groups that naturally reflect scale, opacity, and position without additional computational overhead during rasterization. Second, we introduce a relaxed rigidity constraint within each group. Rather than enforcing identical 3D displacements, we preserve directional consistency and temporal distribution, allowing non-rigid deformations, while preventing artifacts such as geometric inconsistency and floaters caused by incoherent motions.

To demonstrate the versatility of our approach, we integrate our regularization into four representative 4DGS frameworks: RTD [58], Ex4DGS [22], MoDecGS [20], and Grid4D [63]. We evaluate on both the synthetic D-NeRF [41] dataset and real-world benchmarks, HyperNeRF [38] and NeRF-DS [65]. Across all baselines and datasets, our method consistently improves rendering quality, achieving state-of-the-art performance. On the D-NeRF dataset, our approach improves PSNR by an average of 1.19 dB over the baselines. While the proposed method increases training time by around 2 times, it does not alter the underlying model architectures and therefore introduces no additional cost during rendering.

- We propose a framework that enforces locally coherent motion without relying on external priors.
- We introduce a model-agnostic ray-based grouping strategy and relax rigidity constraints, enabling the flexible representation of motion.
- We demonstrate that our approach achieves state-of-the-art performance on both synthetic and real-world benchmarks, including D-NeRF, HyperNeRF, and NeRF-DS.

2 Related Work

2.1 Monocular Video Novel View Synthesis

Novel view synthesis [18, 23, 34] generates unseen views from input images. While extending this task to dynamic scenes relies on multiview videos [1, 3, 8, 25, 30, 44, 48, 53, 64], utilizing monocular videos [7, 37–39] serves as a highly practical alternative. However, monocular inputs lack cross-view motion cues. Consequently, recent approaches [29, 47, 54, 61, 72] employ prior models for optimization. These methods typically leverage optical flow guidance [62, 69], 2D point tracks [16, 60], and depth estimation [11, 40, 43] to guide 4D scene reconstruction. A critical limitation of these approaches is their strong dependency on the accuracy of the underlying priors. When these external models fail, such as in environments with view-dependent artifacts or textureless backgrounds [61, 72], the resulting synthesis quality degrades. Therefore, designing locally coherent motion guidance that operates independently of prior models is essential.

2.2 Dynamic Extension of 3DGS

The computational efficiency of 3DGS [14, 18] drives the extension of the method to dynamic scenes. Dynamic 3DGS methodologies are classified into three categories. The first category utilizes deformation fields [2, 9, 20, 24, 27, 32, 46, 49, 51, 52, 58, 59, 63, 67, 72] to compute temporal offsets for primitive parameters. The second category employs basis functions. These models adopt Fourier series [17], polynomials [10, 26, 55] or splines [15, 22, 36, 68]. The third category

defines primitives in four-dimensional space [6, 66, 71]. With the exception of native 4D formulations, these methods generate 3D Gaussian primitives to process the primitives through the 3DGS rasterizer. However, regulating the motion of Gaussians is highly ill-posed, necessitating spatial constraints or grouping mechanisms to enforce coherent movement. To achieve this, our approach leverages the rasterization process as a grouping function. This design allows our method to integrate into existing dynamic models without requiring structural modifications.

2.3 Motion Constraints for Dynamic Gaussian Splatting

Imposing motion constraints on Gaussian primitives relies on rigidity assumptions [2, 5, 21, 33, 57]. Existing methods maintain geometric structure by preserving distances among KNN groups or by enforcing local rigidity through As-Rigid-As-Possible (ARAP) formulations [12, 24, 28, 50]. However, the KNN approach [2, 12, 24, 28, 33, 57] groups Gaussians based on Euclidean distance, which ignores the properties of Gaussian primitives, including scale and opacity. Furthermore, strict rigidity assumptions fail during topological changes [38], where the underlying distances between Gaussians change. Therefore, we propose a novel ray-based grouping strategy and a relaxed rigidity constraint to capture the properties of Gaussian primitives while remaining highly adaptable to complex, non-rigid physical movements.

3 Method

3.1 Preliminary

3DGS. Our approach builds upon the point-based differentiable rasterization framework of 3DGS [18]. 3DGS represents a 3D scene with anisotropic Gaussians parameterized by position $\boldsymbol{\mu}$, covariance matrix $\boldsymbol{\Sigma}$, opacity $\boldsymbol{o}$, and color $\boldsymbol{c}$. The covariance $\boldsymbol{\Sigma}$ is parameterized via a diagonal scaling matrix $\boldsymbol{S}$ (from a scale vector $\boldsymbol{s}$) and a rotation matrix $\boldsymbol{R}$ (from a quaternion $\boldsymbol{q}$):

$$\boldsymbol{\Sigma} = \boldsymbol{R} \boldsymbol{S} \boldsymbol{S}^\top \boldsymbol{R}^\top. \quad (1)$$

To render an image, 3D Gaussians are projected to the image plane using an approximated graphics pipeline [73]. The 2D covariance $\boldsymbol{\Sigma}'$ in camera coordinates is computed as:

$$\boldsymbol{\Sigma}' = \boldsymbol{J} \boldsymbol{W} \boldsymbol{\Sigma} \boldsymbol{W}^\top \boldsymbol{J}^\top, \quad (2)$$

where $\boldsymbol{J}$ is the Jacobian of the affine approximation for perspective projection, and $\boldsymbol{W}$ is the viewing transformation matrix. By discarding the third row and column of $\boldsymbol{\Sigma}'$, the 3D Gaussian is approximated as a 2D anisotropic Gaussian on the image plane. The density of the projected Gaussian is:

$$\mathcal{G}(\boldsymbol{x}) = e^{-\frac{1}{2}(\boldsymbol{x}-\boldsymbol{\mu})^\top \boldsymbol{\Sigma}'^{-1}(\boldsymbol{x}-\boldsymbol{\mu})}. \quad (3)$$

The final pixel color is computed using point-based α -blending, analogous to volume rendering. In 3DGS, $\alpha_i = \boldsymbol{o}_i \mathcal{G}(\boldsymbol{x})$. For Gaussians sorted along a ray, the accumulated color $\boldsymbol{C}$ is defined as follows:

$$C = \sum_{i=1}^N T_i (1 - e^{-\alpha_i}) \mathbf{c}_i \quad \text{s.t.} \quad T_i = e^{-\sum_{j=1}^{i-1} \alpha_j}, \quad (4)$$

where N is the number of visible Gaussians intersecting the ray, and i is the sorted index by depth. We define the contribution weight as $w_i = T_i(1 - e^{-\alpha_i})$.

Deformation-field Dynamic 3DGS Model. Dynamic Gaussian Splatting extends 3DGS by introducing time-dependent offsets to the canonical Gaussian parameters. These offsets are typically predicted via a deformation network conditioned on time and the canonical attributes. For a canonical Gaussian $G_c = (\boldsymbol{\mu}_c, \mathbf{q}_c, \mathbf{s}_c, \mathbf{o}_c)$, the deformation field outputs the time-varying offsets such that the per-frame primitives become:

$$G_t = (\boldsymbol{\mu}_c + \Delta\boldsymbol{\mu}_t, \mathbf{q}_c + \Delta\mathbf{q}_t, \mathbf{s}_c + \Delta\mathbf{s}_t, \mathbf{o}_c + \Delta\mathbf{o}_t), \quad (5)$$

where $\Delta(\cdot)$ denotes time-dependent offsets.

Motion Basis for Dynamic 3DGS Model. Basis-trajectory approaches [17, 22, 26] parameterize most motions through linear combinations of low-dimensional temporal basis. For example, the position $\boldsymbol{\mu}$ can be represented below:

$$\boldsymbol{\mu}_i(t) = \sum_{\ell=1}^L \mathbf{a}_{i,\ell} \phi_\ell(t), \quad (6)$$

where $\{\phi_\ell\}$ are basis functions and $\mathbf{a}_{i,\ell}$ are learnable coefficients. In our experiments, we validate our approach on both deformation-field (Eq. 5) and basis-trajectory (Eq. 6) formulations to demonstrate its applicability across representative dynamic 3DGS architectures.

3.2 Ray-based Gaussian Grouping

Our method begins by selecting Gaussians along each camera ray. To facilitate object-level motion modeling, we require local groups of primitives that exhibit similar motion. In 3DGS, the number of primitives changes during training due to adaptive density control, which would require frequent re-grouping, making an efficient grouping strategy essential. To achieve this, we leverage the standard rasterization pipeline, which inherently aggregates and sorts Gaussians along each pixel ray. We re-purpose this visibility mechanism to construct groups with minimal overhead throughout training. For each pixel p_j , we define the ray-based group $\mathcal{N}_j$ by selecting Gaussians that actively contribute to rendering:

$$\mathcal{N}_j = \{\mathcal{G}_i \mid w_i > \tau\}, \quad (7)$$

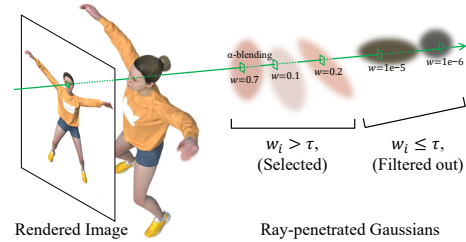


Fig. 2: Ray-based Grouping visualization. We gather the Gaussian penetrated by a ray, and select the Gaussian whose contribution w_i exceeds a threshold τ .

where w_i is the blending weight from Eq. 4, and τ is a threshold. A visualization of grouping is shown in Fig. 2. We subsequently impose local regularizers on these groups to encourage coherent motion while preserving local shape structure.

This design is fundamentally motivated by the physical and geometric properties of the volume rendering process [34]. The blending weight w_i acts as an implicit, **occlusion-aware filter**. As the accumulated transmittance attenuates sharply behind opaque occluders [18], thresholding by τ inherently isolates spatially contiguous Gaussians on the unoccluded surface, thereby preventing the entanglement of foreground and background Gaussians. By leveraging this filtered visibility, grouping Gaussians along view rays establishes a direct bridge between 2D observations and 3D geometry. Instead of relying on explicit 3D distance metrics that might erroneously cluster physically close but structurally independent parts [42], our ray-based grouping establishes a dense set of spatio-temporal constraints to robustly regularize object motion. Furthermore, this formulation naturally yields highly flexible and adaptive group sizes. As demonstrated by the wide distribution in Fig. 3, our approach dynamically accommodates varying local complexities, ranging from thin structures to dense volumes, without the need for explicit structural priors or heuristic guidance.

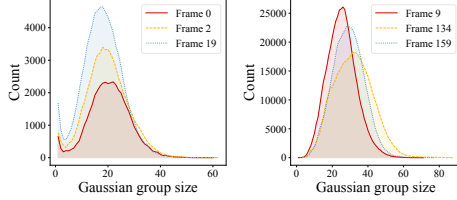


Fig. 3: Group size visualization. Distribution of Gaussian group sizes for the D-NeRF (left) and HyperNeRF (right). The plot shows the counts of groups with varying sizes across different image frames. Groups of size zero are omitted.

3.3 Motion Coherence Regularization

We first introduce a motion coherence regularization (MCR) term that encourages coherent dynamics within each ray-based group $\mathcal{N}_j$ without relying on external supervision such as optical flow or monocular depth [72]. Instead, we regularize directional consistency while allowing magnitudes to adapt. Let $\boldsymbol{\mu}_{i,t}$ be the 3D position of Gaussian $\mathcal{G}_i \in \mathcal{N}_j$ at time t . For a temporal offset Δt , the displacement is:

$$\mathbf{d}_{i,t} = \boldsymbol{\mu}_{i,t+\Delta t} - \boldsymbol{\mu}_{i,t}. \quad (8)$$

We compute the mean displacement of a group:

$$\bar{\mathbf{d}}_{\mathcal{N}_j,t} = \frac{1}{|\mathcal{N}_j|} \sum_{\mathcal{G}_i \in \mathcal{N}_j} \mathbf{d}_{i,t}. \quad (9)$$

We then penalize directional inconsistency via a cosine-similarity loss:

$$\mathcal{L}_{\text{MCR}} = 1 - \frac{\mathbf{d}_{i,t} \cdot \bar{\mathbf{d}}_{\mathcal{N}_j,t}}{\|\mathbf{d}_{i,t}\| \|\bar{\mathbf{d}}_{\mathcal{N}_j,t}\| + \epsilon}, \quad (10)$$

where ϵ is a small positive number. The loss is applied only to moving Gaussians satisfying $\|\mathbf{d}_{i,t}\| > 0.0001$. This regularization encourages Gaussians that contribute to the same pixel to move coherently, while still permitting spatially varying motion magnitudes. Crucially, we do not penalize differences in displacement magnitudes. This relaxed constraint is essential, as enforcing uniform motion magnitudes would erroneously impose strict rigid translations by conflating coherent motion with identical displacement.

3.4 Spectral Regularization

MCR effectively aligns the overall trajectories of Gaussians, but lacks an explicit mechanism to preserve local spatial structures over time. A common geometric prior for this task is local rigidity, typically enforced through ARAP energy functions [12, 50] that penalize changes in pairwise distances and relative orientations. However, enforcing such point-to-point rigidity is overly restrictive for dynamic scenes, as it severely hinders natural non-rigid deformations and smooth motion.

Rather than imposing strict point correspondences, we propose a spectral regularization (SR) that maintains the distributional shape of the group. We compute the covariance matrix of Gaussian positions at time steps t and $t + \Delta t$ and penalize the difference between their eigenvalue spectra. Let K_t and $K_{t+\Delta t}$ be covariances of the same ray group at t and $t + \Delta t$. Let $\sigma_{t,r}$ and $\sigma_{t+\Delta t,r}$ be the r -th eigenvalues of K_t and $K_{t+\Delta t}$ in ascending order, respectively. We define the covariance consistency loss as follows:

$$\mathcal{L}_{\text{SR}} = \sum_{r \in \{1,2,3\}} \text{Huber}(\sigma_{t,r}, \sigma_{t+\Delta t,r}), \quad (11)$$

where $\text{Huber}(\cdot, \cdot)$ is the standard Huber loss [13]. We apply SR only to groups with more than one Gaussian.

Matching eigenvalue spectra preserves local shape statistics over time, while remaining invariant to rigid rotations and allowing flexible non-rigid motions. This preserves the spatial volume while allowing localized deformations. This formulation prevents shape distortion even when the group forms a broad spatial distribution. Since the loss function does not minimize the overall scale of the eigenvalue spectrum, the regularizer does not force the group to contract even if it contains distant Gaussians. Instead, it selectively targets Gaussians that exhibit displacements large enough to alter the object shape. The regularization guides these deviating Gaussians to maintain spatial proximity with their grouped neighbors, restricting only the movements that disrupt the structure.

ARAP baseline for comparison. For completeness, we additionally consider an ARAP-style rigidity prior [12, 50] as a *comparative* alternative to $\mathcal{L}_{\text{SR}}$ in our experiments. Concretely, for each Gaussian i with neighbors $\mathcal{M}(i)$, we estimate a local rotation $\mathbf{R}_{i,t}$ that best aligns relative vectors across time and penalize residual distortion:

$$\mathcal{L}_{\text{ARAP}} = \sum_i \sum_{j \in \mathcal{M}(i)} \|(\boldsymbol{\mu}_{i,t+\Delta t} - \boldsymbol{\mu}_{j,t+\Delta t}) - \mathbf{R}_{i,t}(\boldsymbol{\mu}_{i,t} - \boldsymbol{\mu}_{j,t})\|_2^2. \quad (12)$$

We use this term only to contextualize the effect of strict local rigidity relative to our spectral regularization.

3.5 Welford’s Algorithm

We employ Welford’s algorithm [56] for efficient, online covariance calculation. Consider a group consisting of N Gaussians. Let K_{i-1} represent the covariance matrix computed from the first $i - 1$ Gaussian positions. When incorporating the i -th Gaussian, the covariance K_i is updated recursively as follows:

$$K_i = \frac{i-1}{i} K_{i-1} + \frac{1}{i} P_i, \quad P_i := (\boldsymbol{\mu}_i - \bar{\boldsymbol{\mu}}_{i-1})(\boldsymbol{\mu}_i - \bar{\boldsymbol{\mu}}_{i-1})^\top, \quad i = 1, \dots, N, \quad (13)$$

where $\bar{\boldsymbol{\mu}}_i = \frac{1}{i} \sum_{j=1}^i \boldsymbol{\mu}_j$ is the running mean. By induction,

$$K_N = \sum_{i=1}^N \left(\frac{1}{i} \prod_{k=i+1}^N \frac{k-1}{k} \right) P_i = \frac{1}{N} \sum_{i=1}^N P_i. \quad (14)$$

This enables covariance computation along each ray in a single pass and can be integrated into the rasterization pipeline (a proof in Sec. 10).

3.6 Final Objective

We adopt the photometric losses used in 3DGS, namely the ℓ_1 loss $\mathcal{L}_1$ and the structural dissimilarity loss $\mathcal{L}_{\text{dssim}}$, and augment them with our regularization terms. Our default training objective is:

$$\mathcal{L} = (1 - \lambda_{\text{dssim}}) \mathcal{L}_1 + \lambda_{\text{dssim}} \mathcal{L}_{\text{dssim}} + \lambda_{\text{MCR}} \mathcal{L}_{\text{MCR}} + \lambda_{\text{SR}} \mathcal{L}_{\text{SR}}, \quad (15)$$

where λ_{dssim} , λ_{MCR} , and λ_{SR} are the hyperparameters. In the ARAP comparison setting, we replace $\mathcal{L}_{\text{SR}}$ with $\mathcal{L}_{\text{ARAP}}$ (with its own weight) and set $\lambda_{\text{MCR}} = 0$ while keeping all other components identical.

4 Experiment

4.1 Datasets

We evaluate our method on the D-NeRF [41], HyperNeRF [38], and NeRF-DS [65] datasets. Following prior protocols, we report PSNR, SSIM, and VGG-based LPIPS (LPIPS_V) [70] for D-NeRF; PSNR, MS-SSIM (MS), and AlexNet-based LPIPS (LPIPS_A) for HyperNeRF; and PSNR, SSIM, MS-SSIM, and VGG-based LPIPS for NeRF-DS.

D-NeRF. D-NeRF contains eight synthetic scenes featuring object deformations, captured using a 360° monocular camera setup. Following prior works [41, 58], we perform evaluation at a resolution of 800×800 . We exclude the *Lego* scene because the training and test splits are not temporally aligned [63, 67].

Table 1: Quantitative comparison of our method on the D-NeRF, HyperNeRF and NeRF-DS dataset. We report PSNR, SSIM, and LPIPS_V for D-NeRF, PSNR, MS-SSIM, and LPIPS_A for HyperNeRF, and PSNR, SSIM, MS-SSIM, and LPIPS_V for NeRF-DS. Δ denotes the PSNR improvement when the baseline models are combined with our framework.

Model	D-NeRF				HyperNeRF				NeRF-DS			
	PSNR $\uparrow$	SSIM $\uparrow$	LPIPS _V $\downarrow$	$\Delta\uparrow$	PSNR $\uparrow$	MS-SSIM $\uparrow$	LPIPS _A $\downarrow$	$\Delta\uparrow$	PSNR $\uparrow$	SSIM/MS $\uparrow$	LPIPS _V $\downarrow$	$\Delta\uparrow$
NeRF [34]	18.81	0.881	0.176	-	20.13	0.745	0.424	-	N/A	N/A	N/A	-
D-NeRF [41]	31.69	0.970	0.043	-	N/A	N/A	N/A	-	N/A	N/A	N/A	-
TiNeuVox [7]	33.76	0.980	0.039	-	24.25	0.837	N/A	-	21.61	0.823/-	0.277	-
HyperNeRF [38]	N/A	N/A	N/A	-	22.38	0.814	N/A	-	23.45	- /0.849	0.199	-
NeRF-DS [65]	N/A	N/A	N/A	-	N/A	N/A	N/A	-	23.60	- /0.849	0.182	-
3D-GS [18]	23.39	0.929	0.077	-	20.26	0.657	0.342	-	20.29	0.782/-	0.292	-
D3DGS [67]	40.43	0.992	0.011	-	22.40	N/A	0.275	-	23.61	0.839/-	0.197	-
GaGS [32]	39.07	<u>0.993</u>	0.007	-	24.26	0.793	N/A	-	N/A	N/A	N/A	-
Ex4DGS [22]	29.93	0.968	0.041	-	24.61	0.829	0.199	-	21.73	0.793/0.789	0.251	-
RTD [58]	35.36	0.985	0.022	-	25.17	0.841	0.283	-	21.91	0.812/0.821	0.235	-
MoDec-GS [20]	29.33	0.963	0.048	-	24.92	0.841	<u>0.176</u>	-	22.93	0.822/0.851	0.229	-
Grid4D [63]	42.00	0.994	<u>0.008</u>	-	25.50	0.783	0.185	-	23.41	0.833/0.872	0.196	-
Ex4DGS+Ours	31.04	0.972	0.038	1.11	24.86	0.836	0.233	0.25	21.88	0.798/0.794	0.248	0.15
RTD+Ours	36.46	0.987	0.019	1.10	25.30	0.845	0.273	0.13	22.05	0.813/0.826	0.226	0.14
MoDec-GS+Ours	31.68	0.974	0.035	2.35	25.09	<u>0.846</u>	0.173	0.17	23.76	0.848/0.885	<u>0.184</u>	0.83
Grid4D+Ours	42.20	0.994	<u>0.008</u>	0.20	25.56	0.856	0.198	0.06	<u>23.70</u>	<u>0.840/0.882</u>	0.188	0.32

HyperNeRF. HyperNeRF features real-world scenes including noise, illumination variations, and topological changes. We use the vrig split, which is captured with a handheld dual-camera rig that simultaneously records two views per frame during data collection. The training set is constructed by alternately selecting images from the left and right cameras for each frame, with the remaining images reserved for testing. For all experiments, images are downsampled to half their original resolution, *i.e.*, 536×960 . Camera poses are estimated using COLMAP [45], and point clouds are obtained from RTD [58].

NeRF-DS. NeRF-DS comprises seven real-world scenes with specular objects and is provided at a resolution of 480×270 . The dataset is captured with a dual-camera setup, using the left camera for training and the right-camera for testing.

4.2 Baseline models

We compare our method against four representative baselines: RTD [58], MoDec-GS [20], Grid4D [63], and Ex4DGS [22]. RTD, MoDec-GS, and Grid4D adopt deformation-based formulations, whereas Ex4DGS uses spline parameterization. Specifically, RTD parameterizes the deformation field with HexPlanes [3], MoDec-GS employs scaffold-based representations [31] to model motion, and Grid4D utilizes a hash-grid representation [35] for deformation encoding. In contrast, Ex4DGS represents dynamic Gaussians explicitly through spline trajectories.

Implementation. All experiments are conducted on a single NVIDIA RTX 3090 GPU. We build upon the official implementations of the baseline methods and adopt the dataset-specific parameter settings provided in their respective protocols. For fair comparison, we keep the baseline architectures and hyperparameters unchanged, except for the additional regularization introduced by our method. Detailed implementation specifics are provided in Sec. 8.

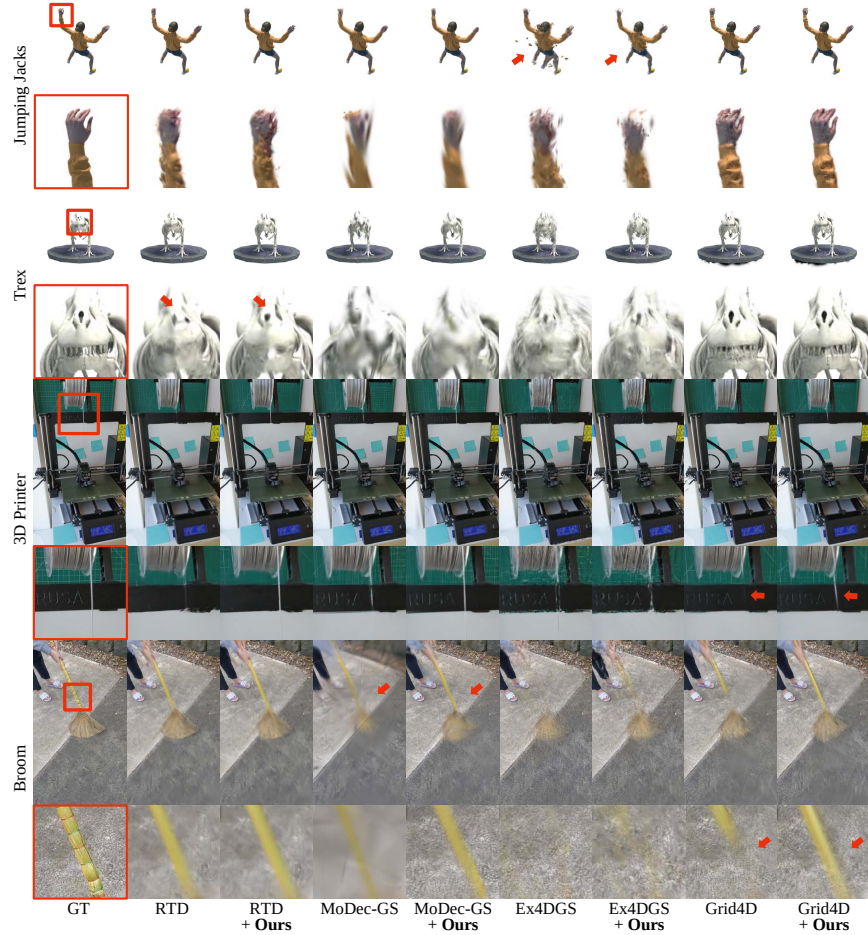


Fig. 4: Qualitative comparisons between baselines and our method on the D-NeRF and HyperNeRF datasets.

4.3 Comparison

Quantitative comparison. Tab. 1 reports quantitative results on the D-NeRF, HyperNeRF, and NeRF-DS. Our method consistently improves reconstruction quality when integrated with all baselines. On the D-NeRF, we observe clear PSNR gains across all models, *e.g.*, +1.11 dB for Ex4DGS and +2.35 dB for MoDec-GS, while Grid4D+Ours achieves the best PSNR of 42.20. On the HyperNeRF, our regularization improves both PSNR and SSIM for most baselines, with Grid4D+Ours obtaining the highest MS-SSIM of 0.856. Notably, on the more challenging NeRF-DS dataset, our method provides larger improvements, particularly for MoDec-GS and Grid4D, leading to significant gains in both PSNR and perceptual quality. These results indicate that our regularization is complementary to existing dynamic 3DGS methods and consistently enhances reconstruction performance across diverse datasets.

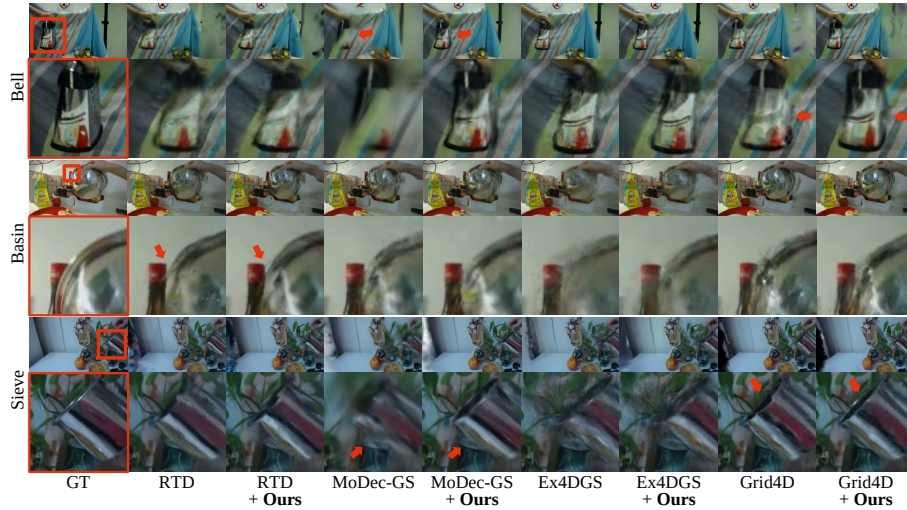


Fig. 5: Qualitative comparisons between baselines and our method on the NeRF-DS dataset.

Qualitative comparison. Figs. 4 and 5 present visual comparisons on the D-NeRF, HyperNeRF, and NeRF-DS datasets. Existing methods produce artifacts such as disappearing objects or distorted shapes, whereas our method preserves the structural integrity of Gaussians over time. On the HyperNeRF, prior methods tend to blur or remove thin structures, like the handle of the broom, while our regularization preserves structures and maintains locally coherent shapes. Similarly, in the *Jumping Jacks* scene of the D-NeRF, our method better retains fine details compared to the baselines. This improvement arises because our regularization aligns Gaussian motions with the underlying object dynamics, preventing incorrect motion updates during image-based optimization. Furthermore, as shown in Fig. 4, our ray-based grouping strategy enables reliable grouping for small structures, like the fingers in *Jumping Jacks* scene and the teeth in *Trex* scene, facilitating accurate shape optimization for methods like Grid4D. On the NeRF-DS *Bell* scene, despite the presence of specular reflections, enforcing physically consistent motion leads to improved reconstruction quality.

Ablation study. Tabs. 2 and 3 and Fig. 6 present ablation studies on the D-NeRF and HyperNeRF datasets using RTD as the base model. We report the average results over all scenes in each dataset. Specifically, we compare our ray-based grouping (RG) with a KNN grouping strategy and analyze the effect of removing each component, including ARAP, MCR, and SR. Except for the introduced components, all hyperparameters remain identical to the baseline configuration. We additionally report the training time to analyze the computational overhead. For the KNN baseline, we follow the strategy used in D3DGS [33], which constructs a group by selecting the 20 nearest Gaussians at every iteration. For ARAP regularization, we adopt the formulation from Sec. 3.4. The

Table 2: Ablation comparison table. Experiments are conducted on the RTD model, and we report the average results on the D-NeRF and HyperNeRF datasets. KNN denotes the use of KNN-based grouping, and ARAP indicates as-rigid-as-possible regularization. RG represents ray-based grouping, MCR denotes motion coherence regularization, and SR denotes spectral regularization. **Full** indicates the configuration where both MCR and SR are applied. Time refers to the training time.

Model	D-NeRF				HyperNeRF			
	PSNR $\uparrow$	SSIM $\uparrow$	LPIPS $\downarrow$	Time(s)	PSNR $\uparrow$	MS-SSIM $\uparrow$	LPIPS $\downarrow$	Time(s)
RTD [58]	35.36	0.9845	0.0224	841	25.17	0.8413	0.2827	1963
RTD+KNN&ARAP	34.80	0.9828	0.0258	<u>1484</u>	23.18	0.7532	0.4606	5808
RTD+KNN& Full	33.29	0.9771	0.0290	2044	25.11	0.8390	0.2856	9586
RTD+RG&ARAP	35.27	0.9844	0.0216	2125	<u>25.25</u>	0.8430	<u>0.2710</u>	13096
RTD+RG&MCR	<u>36.23</u>	<u>0.9866</u>	0.0190	1649	24.23	<u>0.8449</u>	0.2713	<u>4454</u>
RTD+RG&SR	34.98	0.9841	0.0220	1501	24.91	0.8331	0.2961	6694
RTD+RG& Full	36.46	0.9869	<u>0.0193</u>	1924	25.30	0.8564	0.2488	7138

Full configuration denotes the model with both MCR and SR enabled. Detailed numerical results are provided in Sec. 9.

KNN comparison. We first compare our ray-based grouping with the KNN grouping strategy. Since KNN grouping does not account for Gaussian scale or orientation, it often fails to select physically coherent groups. As shown in Tab. 2, this leads to degraded performance for both ARAP and **Full** configurations. Incorrect grouping causes unrelated Gaussians to move together, resulting in inconsistent motion and reduced reconstruction quality.

The qualitative results illustrate this issue. As shown in Fig. 6, KNN grouping frequently produces mis-grouped Gaussians. In the *Hell Warrior* scene, applying strong rigidity constraints such as ARAP to incorrectly grouped Gaussians can even cause parts of the object to disappear. In contrast, our ray-based grouping preserves the structural consistency of the object. Moreover, the visibility-based thresholding prevents grouping across occluded regions and helps preserve fine details, such as the eyes of the chicken in the RG+**Full** configuration.

ARAP comparison. We further evaluate ARAP regularization with both KNN and RG grouping strategies, as reported in Tab. 2. When combined with RG, ARAP produces comparable or improved results compared to KNN. However, on the D-NeRF dataset, the performance still decreases. This is because the rigid transformation assumption in ARAP restricts fine-grained shape variations, preventing the model from capturing detailed geometric changes. We also analyze the individual effects of MCR and SR. MCR encourages directional consistency among neighboring Gaussians but does not explicitly preserve shape. Thus, it works well for scenes with relatively linear motion, such as those in the D-NeRF dataset, but becomes less effective in the more complex dynamics of HyperNeRF. In contrast, SR focuses solely on shape preservation and does not constrain motion, which makes it insufficient when applied alone. These results indicate that both a well-defined grouping strategy and an appropriate motion prior are necessary to achieve consistent performance improvements.

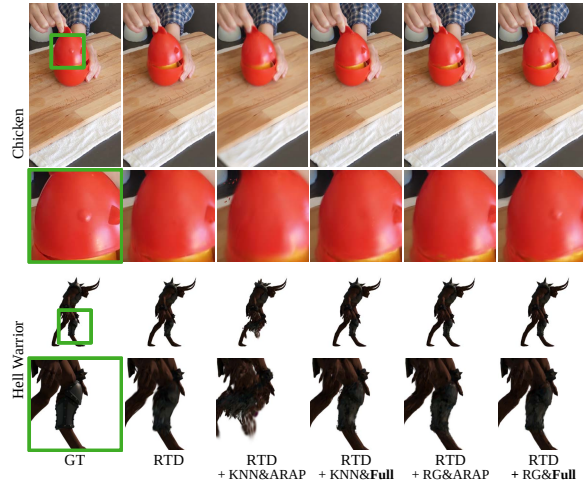


Fig. 6: Ablation results on the RTD model. We compare KNN grouping and ray-based grouping (RG) under different regularization settings, including ARAP, MCR, and SR. **Full** denotes the configuration where both MCR and SR are applied.

Training time. The training time overhead introduced by the proposed regularization is reported in Tab. 2. The increase is more noticeable on the HyperNeRF dataset, which contains a larger number of Gaussians and higher pixel counts. Most of the additional cost arises during rasterization, particularly from covariance processing and SVD operations required for the regularization terms. Despite this overhead, RG remains 6%–25% faster than KNN in training time under the **Full** configuration.

Table 3: Hyperparameter analysis on the Trex scene of the D-NeRF dataset using RTD baseline model. We compare the threshold τ , the MCR and SR loss weights, and the temporal interval. We report PSNR, SSIM, and LPIPS (VGG).

	Temporal Interval			λ_{MCR}			λ_{SR}			τ		
	10	20	30	1e-3	1e-2	1e-1	1e-4	1e-3	1e-2	1e-4	1e-3	1e-1
PSNR	34.40	35.62	34.79	35.62	35.40	35.32	35.29	35.62	35.40	35.62	35.59	34.79
SSIM	0.986	0.987	0.987	0.986	0.987	0.981	0.987	0.988	0.987	0.988	0.988	0.986
LPIPS _V	0.022	0.021	0.022	0.022	0.021	0.028	0.022	0.020	0.020	0.020	0.021	0.022

Hyperparameter analysis. Tab. 3 reports the effects of the visibility threshold τ , the regularization weights, and the temporal interval. The results show that too short temporal intervals provide only weak constraints, while excessively large MCR weights over-constrain non-rigid motion. In contrast, SR is relatively less affected by its weight. Large values of τ also degrade performance by discarding Gaussians that make meaningful contributions to the grouping.

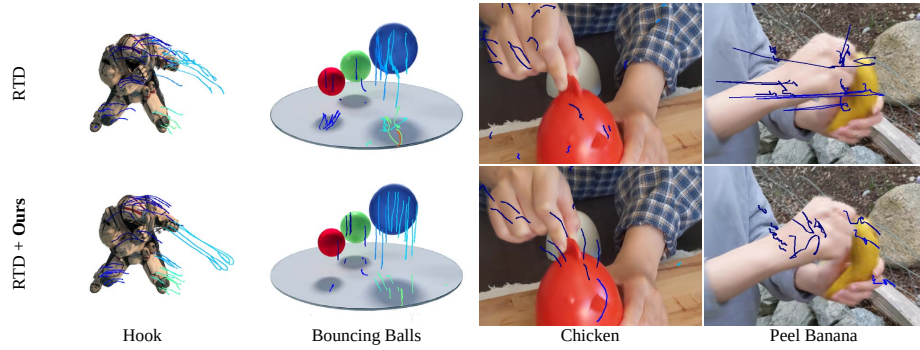


Fig. 7: Gaussian trajectory visualization results of the RTD with our method on the D-NeRF and HyperNeRF datasets. We project the temporal positions of sampled Gaussians onto a fixed camera view.

Trajectory visualization. Fig. 7 visualizes Gaussian trajectories of the RTD model for both the baseline and our method on the D-NeRF and HyperNeRF datasets. We visualize the trajectories of Gaussian positions by projecting them onto a fixed camera to analyze their motion. On the D-NeRF, our method produces trajectories that are better aligned with the underlying physical motion compared to the baseline. While the baseline often exhibits permuted or irregular Gaussian paths, our method results in more structured and temporally consistent trajectories. This behavior arises because our regularization encourages neighboring Gaussians to move coherently. On the HyperNeRF dataset, particularly around the hand region, the baseline often produces Gaussians that drift away from the object surface. In contrast, our method keeps Gaussians consistently attached to the same object region. This result indicates that our approach effectively regularizes outlier Gaussians and improves motion consistency.

5 Conclusion

Modeling locally coherent motion in dynamic 3D Gaussian Splatting remains challenging, as many existing approaches rely heavily on external priors to mitigate incoherent motion. In this work, we propose a method that enables dynamic 3DGS models to learn locally coherent motion directly from image supervision. By introducing a ray-based grouping strategy coupled with motion coherence and spectral regularization, our approach effectively enforces relaxed rigidity constraints. We integrate our approach into four baseline models and evaluate them across three datasets, consistently observing performance improvements. Notably, integrating our method with Grid4D yields state-of-the-art results, demonstrating both the effectiveness and generality of our framework. Our findings confirm that physically grounded motion constraints can significantly improve dynamic 3DGS reconstruction without relying on external priors. Future work will explore finer-grained physical constraints to facilitate robust modeling under highly complex or restricted motion conditions.

6 Acknowledgment

This work was supported in part by the Yonsei University Future-Leading Research Initiative of 2025 (2025-22-0409), the Electronics and Telecommunications Research Institute (ETRI) grant funded by the Korean government (26CC1100, Development of User-Customized Freeform Future Display), and the National Research Foundation of Korea (NRF) grant funded by the Korea government (MSIT) (RS-2024-00338439).

References

1. Attal, B., Huang, J.B., Richardt, C., Zollhoefer, M., Kopf, J., O’Toole, M., Kim, C.: HyperReel: High-fidelity 6-dof video with ray-conditioned sampling. In: CVPR (2023)
2. Bae, J., Kim, S., Yun, Y., Lee, H., Bang, G., Uh, Y.: Per-gaussian embedding-based deformation for deformable 3d gaussian splatting. In: ECCV (2024)
3. Cao, A., Johnson, J.: Hexplane: A fast representation for dynamic scenes. In: CVPR (2023)
4. Chen, Q., Qu, D., Liu, J., Tang, Y., Song, H., Wang, D., Zhao, B., Li, X.: Free-gaussian: Annotation-free control of articulated objects via 3d gaussian splats with flow derivatives. In: AAAI (2025)
5. Das, D., Wewer, C., Yunus, R., Ilg, E., Lenssen, J.E.: Neural parametric gaussians for monocular non-rigid object reconstruction. In: CVPR (2024)
6. Duan, Y., Wei, F., Dai, Q., He, Y., Chen, W., Chen, B.: 4d-rotor gaussian splatting: towards efficient novel view synthesis for dynamic scenes. In: SIGGRAPH (2024)
7. Fang, J., Yi, T., Wang, X., Xie, L., Zhang, X., Liu, W., Nießner, M., Tian, Q.: Fast dynamic radiance fields with time-aware neural voxels. In: SIGGRAPH Asia 2022 Conference Papers (2022)
8. Fridovich-Keil, S., Meanti, G., Warburg, F.R., Recht, B., Kanazawa, A.: K-Planes: Explicit radiance fields in space, time, and appearance. In: CVPR (2023)
9. Gao, Q., Meng, J., Wen, C., Chen, J., Zhang, J.: HiCoM: Hierarchical coherent motion for dynamic streamable scenes with 3d gaussian splatting. NeurIPS (2024)
10. Hu, B., Li, Y., Xie, R., Xu, B., Dong, H., Yao, J., Lee, G.H.: Learnable infinite taylor gaussian for dynamic view rendering. In: CVPR (2025)
11. Hu, W., Gao, X., Li, X., Zhao, S., Cun, X., Zhang, Y., Quan, L., Shan, Y.: DepthCrafter: Generating consistent long depth sequences for open-world videos. In: CVPR (2025)
12. Huang, Y.H., Sun, Y.T., Yang, Z., Lyu, X., Cao, Y.P., Qi, X.: SC-GS: Sparse-controlled gaussian splatting for editable dynamic scenes. In: CVPR (2024)
13. Huber, P.J.: Robust Estimation of a Location Parameter. *The Annals of Mathematical Statistics* (1964)
14. Hyung, J., Hong, S., Hwang, S., Lee, J., Choo, J., Kim, J.H.: Effective rank analysis and regularization for enhanced 3d gaussian splatting. In: Conference on Neural Information Processing Systems (NeurIPS). NeurIPS (2024)
15. Kang, T., Park, J., Lee, K., Lee, Y.: Clustered error correction with grouped 4d gaussian splatting. In: SIGGRAPH Asia 2025 Conference Papers (2025)
16. Karaev, N., Rocco, I., Graham, B., Neverova, N., Vedaldi, A., Rupprecht, C.: CoTracker: It is better to track together. In: ECCV (2024)

17. Katsumata, K., Vo, D.M., Nakayama, H.: A compact dynamic 3d gaussian representation for real-time dynamic view synthesis. In: ECCV (2024)
18. Kerbl, B., Kopanas, G., Leimkühler, T., Drettakis, G.: 3d gaussian splatting for real-time radiance field rendering. ACM TOG (2023)
19. Kratimenos, A., Lei, J., Daniilidis, K.: DynMF: Neural motion factorization for real-time dynamic view synthesis with 3d gaussian splatting. In: ECCV (2024)
20. Kwak, S., Kim, J., Jeong, J.Y., Cheong, W.S., Oh, J., Kim, M.: MoDec-GS: Global-to-local motion decomposition and temporal interval adjustment for compact dynamic 3d gaussian splatting. In: CVPR (2025)
21. Lee, J., Choi, H., Cho, M.: Space-time forecasting of dynamic scenes with motion-aware gaussian grouping. CVPR (2026)
22. Lee, J., Jung, H., Bae, I., Won, C., Jeon, H.G.: Fully explicit dynamic gaussian splatting. In: NeurIPS (2024)
23. Lee, J., Jung, H., Park, J.H., Bae, I., Jeon, H.G.: Geometry-aware projective mapping for unbounded neural radiance fields. In: ICLR (2024)
24. Li, D., Huang, S.S., Lu, Z., Duan, X., Huang, H.: ST-4DGS: Spatial-temporally consistent 4d gaussian splatting for efficient dynamic scene rendering. In: ACM SIGGRAPH 2024 Conference Papers. Association for Computing Machinery, New York, NY, USA (2024)
25. Li, T., Slavcheva, M., Zollhoefer, M., Green, S., Lassner, C., Kim, C., Schmidt, T., Lovegrove, S., Goesele, M., Newcombe, R., et al.: Neural 3d video synthesis from multi-view video. In: CVPR (2022)
26. Li, Z., Chen, Z., Li, Z., Xu, Y.: Spacetime gaussian feature splatting for real-time dynamic view synthesis. In: CVPR (2024)
27. Liang, Y., Xu, T., Kikuchi, Y.: Himor: Monocular deformable gaussian reconstruction with hierarchical motion representation. In: CVPR (2025)
28. Lin, Y., Dai, Z., Zhu, S., Yao, Y.: Gaussian-flow: 4d reconstruction with dynamic 3d gaussian particle. In: CVPR (2024)
29. LIU, Q., Liu, Y., Wang, J., Lyu, X., Wang, P., Wang, W., Hou, J.: MoDGS: Dynamic gaussian splatting from casually-captured monocular videos with depth priors. In: ICLR (2025)
30. Lombardi, S., Simon, T., Saragih, J., Schwartz, G., Lehrmann, A., Sheikh, Y.: Neural volumes: learning dynamic renderable volumes from images. ACM TOG (2019)
31. Lu, T., Yu, M., Xu, L., Xiangli, Y., Wang, L., Lin, D., Dai, B.: Scaffold-gs: Structured 3d gaussians for view-adaptive rendering. In: CVPR (2024)
32. Lu, Z., Guo, X., Hui, L., Chen, T., Yang, M., Tang, X., Zhu, F., Dai, Y.: 3d geometry-aware deformable gaussian splatting for dynamic view synthesis. In: CVPR (2024)
33. Luiten, J., Kopanas, G., Leibe, B., Ramanan, D.: Dynamic 3d gaussians: Tracking by persistent dynamic view synthesis. In: 3DV (2024)
34. Mildenhall, B., Srinivasan, P., Tancik, M., Barron, J., Ramamoorthi, R., Ng, R.: Nerf: Representing scenes as neural radiance fields for view synthesis. In: ECCV (2020)
35. Müller, T., Evans, A., Schied, C., Keller, A.: Instant neural graphics primitives with a multiresolution hash encoding. ACM Trans. Graph. **41**(4) (2022)
36. Park, J., Bui, M.Q.V., Bello, J.L.G., Moon, J., Oh, J., Kim, M.: SplineGS: Robust motion-adaptive spline for real-time dynamic 3d gaussians from monocular video. In: CVPR (2025)
37. Park, K., Sinha, U., Barron, J.T., Bouaziz, S., Goldman, D.B., Seitz, S.M., Martin-Brualla, R.: Nerfies: Deformable neural radiance fields. In: CVPR (2021)

38. Park, K., Sinha, U., Hedman, P., Barron, J.T., Bouaziz, S., Goldman, D.B., Martin-Brualla, R., Seitz, S.M.: Hypernerf: a higher-dimensional representation for topologically varying neural radiance fields. *ACM TOG* (2021)
39. Park, S., Son, M., Jang, S., Ahn, Y.C., Kim, J.Y., Kang, N.: Temporal interpolation is all you need for dynamic neural radiance fields. In: *CVPR* (2023)
40. Piccinelli, L., Yang, Y.H., Sakaridis, C., Segu, M., Li, S., Van Gool, L., Yu, F.: UniDepth: Universal monocular metric depth estimation. In: *CVPR* (2024)
41. Pumarola, A., Corona, E., Pons-Moll, G., Moreno-Noguer, F.: D-NeRF: Neural Radiance Fields for Dynamic Scenes. In: *CVPR* (2021)
42. Qi, C.R., Yi, L., Su, H., Guibas, L.J.: Pointnet++: Deep hierarchical feature learning on point sets in a metric space. *NeurIPS* (2017)
43. Ranftl, R., Lasinger, K., Hafner, D., Schindler, K., Koltun, V.: Towards robust monocular depth estimation: Mixing datasets for zero-shot cross-dataset transfer. *IEEE TPAMI* **44**(3) (2022)
44. Sabater, N., Boisson, G., Vandame, B., Kerbiriou, P., Babon, F., Hog, M., Langlois, T., Gendrot, R., Bureller, O., Schubert, A., Allie, V.: Dataset and pipeline for multi-view light-field video. In: *CVPRW* (2017)
45. Schonberger, J.L., Frahm, J.M.: Structure-from-motion revisited. In: *CVPR* (2016)
46. Shaw, R., Nazarczuk, M., Song, J., Moreau, A., Catley-Chandar, S., Dhano, H., Pérez-Pellitero, E.: SWinGS: sliding windows for dynamic 3d gaussian splatting. In: *ECCV* (2024)
47. Shin, J., Shaw, R., Shin, S., Zhang, Z., Jeon, H.G., Perez-Pellitero, E.: Chroma: Consistent harmonization of multi-view appearance via bilateral grid prediction. *arXiv preprint arXiv:2507.15748* (2025)
48. Song, L., Chen, A., Li, Z., Chen, Z., Chen, L., Yuan, J., Xu, Y., Geiger, A.: NeRF-Player: A streamable dynamic scene representation with decomposed neural radiance fields. *IEEE TVCG* (2023)
49. Song, Y.J., Kwon, K., Lee, J.L., Kim, J.H., Zhang, B.T.: Kinematics-driven gaussian shape deformation for blurry monocular dynamic scenes. In: *ICML* (2026)
50. Sorkine, O., Alexa, M.: As-rigid-as-possible surface modeling. In: *SGP* (2007)
51. Sun, J., Jiao, H., Li, G., Zhang, Z., Zhao, L., Xing, W.: 3DGStream: On-the-fly training of 3d gaussians for efficient streaming of photo-realistic free-viewpoint videos. In: *CVPR* (2024)
52. Wan, D., Lu, R., Zeng, G.: Superpoint gaussian splatting for real-time high-fidelity dynamic scene reconstruction. In: *ICML* (2024)
53. Wang, F., Tan, S., Li, X., Tian, Z., Song, Y., Liu, H.: Mixed neural voxels for fast multi-view video synthesis. In: *ICCV* (2023)
54. Wang, Q., Ye, V., Gao, H., Zeng, W., Austin, J., Li, Z., Kanazawa, A.: Shape of Motion: 4d reconstruction from a single video. In: *ICCV* (2025)
55. Wang, Y., Yang, P., Xu, Z., Sun, J., Zhang, Z., Chen, Y., Bao, H., Peng, S., Zhou, X.: FreeTimeGS: Free gaussian primitives at anytime anywhere for dynamic scene reconstruction. In: *CVPR* (2025)
56. Welford, B.P.: Note on a method for calculating corrected sums of squares and products. *Technometrics* **4**(3), 419–420 (1962)
57. Wu, D., Liu, F., Hung, Y.H., Qian, Y., Zhan, X., Duan, Y.: 4D-Fly: Fast 4d reconstruction from a single monocular video. In: *CVPR* (2025)
58. Wu, G., Yi, T., Fang, J., Xie, L., Zhang, X., Wei, W., Liu, W., Tian, Q., Xinggang, W.: 4d gaussian splatting for real-time dynamic scene rendering. In: *CVPR* (2024)
59. Wu, J., Peng, R., Wang, Z., Xiao, L., Tang, L., Yan, J., Xiong, K., Wang, R.: Swift4D: Adaptive divide-and-conquer gaussian splatting for compact and efficient reconstruction of dynamic scene. In: *ICLR* (2025)

60. Xiao, Y., Wang, Q., Zhang, S., Xue, N., Peng, S., Shen, Y., Zhou, X.: Spatial-Tracker: Tracking any 2d pixels in 3d space. In: CVPR (2024)
61. Xie, L., Julin, J., Niinuma, K., Jeni, L.A.: Gaussian splatting lucas-kanade. In: ICLR (2025)
62. Xu, H., Zhang, J., Cai, J., Rezatofighi, H., Tao, D.: GMFlow: Learning optical flow via global matching. In: CVPR (2022)
63. Xu, J., Fan, Z., Yang, J., Xie, J.: Grid4D: 4D decomposed hash encoding for high-fidelity dynamic gaussian splatting. NeurIPS (2024)
64. Xu, Z., Peng, S., Lin, H., He, G., Sun, J., Shen, Y., Bao, H., Zhou, X.: 4k4d: Real-time 4d view synthesis at 4k resolution. In: CVPR (2024)
65. Yan, Z., Li, C., Lee, G.H.: Nerf-ds: Neural radiance fields for dynamic specular objects. In: CVPR (2023)
66. Yang, Z., Yang, H., Pan, Z., Zhang, L.: Real-time photorealistic dynamic scene representation and rendering with 4d gaussian splatting. In: ICLR (2023)
67. Yang, Z., Gao, X., Zhou, W., Jiao, S., Zhang, Y., Jin, X.: Deformable 3d gaussians for high-fidelity monocular dynamic scene reconstruction. In: CVPR (2024)
68. Yoon, J., Han, S., Oh, J., Lee, M.: SplineGS: Learning smooth trajectories in gaussian splatting for dynamic scene reconstruction. In: ICLR (2025)
69. Zachary Teed, J.D.: RAFT: Recurrent all-pairs field transforms for optical flow. In: ECCV (2024)
70. Zhang, R., Isola, P., Efros, A.A., Shechtman, E., Wang, O.: The unreasonable effectiveness of deep features as a perceptual metric. In: CVPR (2018)
71. Zhang, X., Liu, Z., Zhang, Y., Ge, X., He, D., Xu, T., Wang, Y., Lin, Z., Yan, S., Zhang, J.: Mega: Memory-efficient 4d gaussian splatting for dynamic scenes. In: ICCV (2025)
72. Zhu, R., Liang, Y., Chang, H., Deng, J., Lu, J., Yang, W., Zhang, T., Zhang, Y.: MotionGS: Exploring explicit motion guidance for deformable 3d gaussian splatting. In: NeurIPS (2024)
73. Zwicker, M., Pfister, H., Van Baar, J., Gross, M.: EWA volume splatting. In: VIS (2001)