
LangSplatV2: High-dimensional 3D Language Gaussian Splatting with 450+ FPS

Wanhua Li^{1,*} Yujie Zhao^{1,2,*} Minghan Qin^{3,*} Yang Liu³ Yuanhao Cai⁴
Chuang Gan^{5,6} Hanspeter Pfister¹

¹Harvard University ²University of Chinese Academy of Sciences ³Tsinghua University
⁴Johns Hopkins University ⁵MIT-IBM Watson AI Lab ⁶UMass Amherst

Abstract

In this paper, we introduce LangSplatV2, which achieves high-dimensional feature splatting at 476.2 FPS and 3D open-vocabulary text querying at 384.6 FPS for high-resolution images, providing a $42 \times$ speedup and a $47 \times$ boost over LangSplat respectively, along with improved query accuracy. LangSplat employs Gaussian Splatting to embed 2D CLIP language features into 3D, significantly enhancing speed and learning a precise 3D language field with SAM semantics. Such advancements in 3D language fields are crucial for applications that require language interaction within complex scenes. However, LangSplat does not yet achieve real-time inference performance (8.2 FPS), even with advanced A100 GPUs, severely limiting its broader application. In this paper, we first conduct a detailed time analysis of LangSplat, identifying the heavyweight decoder as the primary speed bottleneck. Our solution, LangSplatV2 assumes that each Gaussian acts as a sparse code within a global dictionary, leading to the learning of a 3D sparse coefficient field that entirely eliminates the need for a heavyweight decoder. By leveraging this sparsity, we further propose an efficient sparse coefficient splatting method with CUDA optimization, rendering high-dimensional feature maps at high quality while incurring only the time cost of splatting an ultra-low-dimensional feature. Our experimental results demonstrate that LangSplatV2 not only achieves better or competitive query accuracy but is also significantly faster. Codes and demos are available at our project page: <https://langsplat-v2.github.io>.

1 Introduction

Seamless and intuitive interactions between humans and complex 3D environments [1, 2] are paramount for a wide range of applications, such as augmented reality [3, 4] and intelligent robotics [5, 6, 7]. To achieve such interactions, systems must understand and respond to natural language queries in real time. This capability hinges on advancements in 3D language fields—a technology at the intersection of vision-language models [8, 9] and 3D environmental modeling [10, 11].

LangSplat [12], a pioneering model in this domain, leverages 3D Gaussian Splatting to embed 2D CLIP language features into 3D spaces. This method has significantly accelerated the 3D query time, achieving speeds up to $199 \times$ faster than its predecessors [13], which is critical for applications in scenarios where rapid response is essential. Although many recent improvements have been proposed [14, 15], the inference speed remains a major concern, especially at high resolutions, which hinders their widespread application.

Real-time querying is crucial for applications such as real-time navigation [6], interactive gaming [16], and on-the-fly educational tools [17], where delays can disrupt user experience and functionality. To

* Equal contribution.

identify the speed bottleneck of LangSplat, we decompose the entire querying process into three stages: rendering, decoding, and post-processing. Then, we provide a detailed time analysis on the LERF dataset with one A100 GPU. Results in Table 1 show that each query takes 122.1 ms for LangSplat. With some simple engineering modifications, the rendering and post-processing time can be significantly reduced, we term this new version of LangSplat as LangSplat*. The performance of LangSplat* revealed that the decoding stage, which is responsible for transforming low-dimensional latent features into high-dimensional CLIP features with a heavy-weight multi-layer perceptron (MLP), is the primary bottleneck. LangSplat introduces the MLP decoder to significantly reduce the training memory and time cost. Consequently, an additional decoding stage with an MLP is inevitable for test-time querying and reducing the dimensionality of high-dimensional features lowers the accuracy of the queries. However, simply removing the decoder and directly training the high-dimensional CLIP feature for each Gaussian dramatically increases the rendering time. As shown in Figure 1, as the feature splatting dimension increases, the rendering time of LangSplat significantly increases. Specifically, rendering a feature with a dimension of 1536 (assuming we render three semantic levels of 512-dimensional CLIP feature fields in parallel) is 15 times slower than rendering a 3-dimensional field on one A100 GPU.

To address the decoding bottleneck in LangSplat, we introduce LangSplatV2, which drastically enhances querying speed without sacrificing querying accuracy. As shown in Figure 1, our LangSplatV2 successfully decouples rendering speed from the dimensionality of rendering features, enabling the rendering of high-dimensional features at the computational cost of splatting an ultra-low-dimensional feature. As the millions of 3D Gaussian points actually represent a limited number of unique semantics for a 3D scene, our LangSplatV2 assumes that each Gaussian can be represented as a sparse coding over a group of global basis vectors. Then, we derive that learning a high-dimensional feature field is equivalent to learning a 3D coefficient field and a global codebook. For each 3D Gaussian point, instead of augmenting a language feature, we learn sparse coefficients. During testing, we utilize the sparsity and propose an efficient sparse coefficient splatting method with CUDA optimization, which effectively decouples the rendering dimension with feature dimensions. As our method removes the autoencoder and directly learns the 3D field from high-dimensional 2D CLIP features, more accurate language fields are also obtained. Experimental results demonstrate that our LangSplatV2 not only achieves higher accuracy but also delivers a substantial speedup: $42 \times$ faster than LangSplat for high-dimensional feature rendering and $47 \times$ faster for open-vocabulary 3D querying.

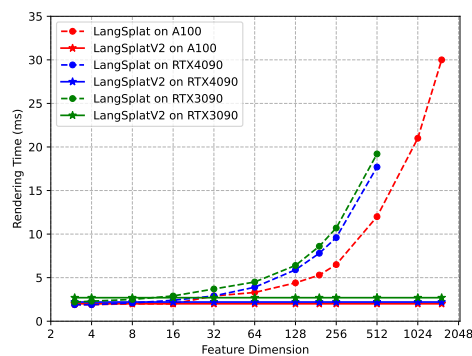


Figure 1: Feature rendering time comparison with different GPUs. Note that the less advanced GPUs (RTX 3090 and RTX 4090) cannot accommodate the LangSplat model with feature dimensions of 1024 or higher due to running out of memory.

2 Related Work

3D Gaussian Splatting. Recently, 3D Gaussian Splatting (3D-GS) [11] has received huge attention [18, 19] due to its speed advantage [20, 21] over neural radiance fields (NeRFs) [10]. As a fundamental 3D modeling technology, it's been widely used in many downstream tasks [22, 23, 24, 25, 26, 27, 28, 29]. Yang *et al.* [30] extended 3D Gaussian Splatting to dynamic scenes by modeling deformable 3D Gaussians. Concurrently, 4D Gaussian Splatting [31] proposes to use multi-resolution voxel planes to model deformed Gaussians and attains high-quality video rendering results in real-time. Due to the explicit 3D modeling nature of 3D Gaussian Splatting, it has also been used for 3D editing. Gaussianeditor [32] offers swift and controllable 3D editing by tracing the editing target throughout the training process. Another concurrent work [33] edits 3D scenes using text instruction, which attains two times faster training speed compared with Instruct-NeRF2NeRF [34]. Many methods [35, 36, 37] utilize 3D Gaussian Splatting for text-driven 3D generation. GSGEN [38] incorporates direct geometric priors from 3D Gaussians and obtains a text-to-3D generation model that

captures high-frequency components. 3D-GS has also been adopted for generative 3D reconstruction, which has recently attracted significant interest owing to its effectiveness with sparse observations. The first framework of its kind was introduced by Liu *et al.* [39], who pioneered generative 3D reconstruction by modeling it as a temporal generation task. Wang *et al.* [40] subsequently accelerated this approach by using one-step diffusion. Recently, 3D language fields have also embraced 3D-GS. LangSplat [12] adopts 3D-GS to model a 3D language field and is $199 \times$ faster than LERF [13]. However, LangSplat is still not a real-time method, which undermines its broad application prospects.

3D Language Field. The concept of 3D language fields [41, 42, 43, 44, 45, 46] has emerged as an intersection of vision-language models and 3D modeling, aiming to create interactive environments that can be manipulated and queried using natural language. LERF [13] explores 3D open-vocabulary text queries by embedding the CLIP feature into a 3D field. Liu *et al.* [47] also distilling knowledge from pre-trained foundation models like CLIP [8] and DINO [48] into NeRF [10]. It further proposes relevancy-distribution alignment and feature-distribution alignment losses to address the CLIP feature ambiguity issue. LangSplat improves LERF with 3D Gaussian Splatting and attains a significant speedup. There are also some other works [49, 50, 51] employing 3D Gaussian Splatting for 3D language field modeling, they usually suffer from imprecise language fields like LERF, which has been well addressed by LangSplat. Following LangSplat, GOI [52] proposes simplifying feature selection by dividing the feature space with a hyperplane, keeping only the features most relevant to the query. GAGS [15] enhances multiview consistency by linking SAM’s prompt point density with camera distances and introducing an unsupervised granularity factor to selectively distill consistent 2D features. However, these methods both use a decoder to map the low-dimensional feature to high-dimensional space, which is a primary bottleneck of real-time query as shown in section 3.2.

3D Gaussian Splatting Acceleration. To reduce computational overhead and improve rendering speed, recent works have proposed various strategies to accelerate 3D-GS. One line of work compresses the Gaussian representation using vector quantization and codebooks. For instance, Compressed 3D Gaussian Splatting [53] adopts sensitivity-aware clustering and quantization-aware training to compress both directional colors and Gaussian parameters, achieving up to $31 \times$ compression and $4 \times$ rendering speedup. CompGS [54] similarly applies K-means-based quantization to achieve over $40 \times$ storage reduction with minimal quality loss. Other approaches focus on structural regularization or architectural simplification. Self-Organizing Gaussian Grids [55] reorganize Gaussian parameters into spatially coherent 2D grids, reducing redundancy via local homogeneity. LightGaussian [56] uses pruning based on global significance and vector quantization, achieving $15 \times$ compression and 200+ FPS rendering. C3DGS [57] proposes to reduce the number of Gaussians via a learnable mask and introduces compact view-dependent color encoding, achieving a $25 \times$ reduction in storage. Techniques like Speedy-Splat [58] and SORT-Free Splatting [59] improve runtime performance through optimized rasterization and differentiable approximations of alpha blending. In contrast to these works that focus on accelerating RGB-based scene rendering, our method tackles the unique challenge of handling high-dimensional Gaussian features. Extending prior compression techniques [53, 56] to our setting would require learning and quantizing such high-dimensional attributes per Gaussian, leading to excessive memory usage. Moreover, directly rendering with full 512D features remains computationally demanding. In contrast, our method circumvents the need to directly learn and render high-dimensional features for 3D Gaussian points.

3 Proposed Approach

3.1 Preliminaries

3D Gaussian Splatting employs point-based rendering technologies [60, 61] and models the scene geometry as a set of 3D Gaussian points. Each Gaussian point is associated with the following attributes: Gaussian center position $\mu \in \mathbb{R}^3$, a covariance matrix $\Sigma \in \mathbb{R}^{3 \times 3}$, color described $c \in \mathbb{R}^3$ by spherical harmonic (SH) coefficients, and opacity $o \in \mathbb{R}$. Each Gaussian point is represented as:

$$G(x) = \exp\left(-\frac{1}{2}(x - \mu)^\top \Sigma^{-1}(x - \mu)\right). \quad (1)$$

The covariance matrix Σ is further represented with a rotation matrix and a scaling factor matrix:

$$\Sigma = RSS^\top R^\top, \quad (2)$$

where $R \in \mathbb{R}^4$ denotes the rotation matrix and $S \in \mathbb{R}^3$ represents the scaling matrix.

Table 1: The stage-wise time cost (ms) for LangSplat and our improvements on the LERF dataset with one A100 GPU. LangSplat* is modified from LangSplat with simple engineering optimization. LangSplatV2 achieves a speed of up to 384.6 FPS for open-vocabulary 3D querying.

Method	Rendering	Decoding	Post-Processing	Total	Speed (FPS)
LangSplat	6.0	83.1	33.0	122.1	8.2
LangSplat*	2.0	83.1	0.5	85.6	11.7
LangSplatV2	2.0	0.1	0.5	2.6	384.6

Based on EWA volume splatting [62], 3D Gaussian Splatting projects the 3D Gaussian points onto a 2D image plane, which blends the ordered Gaussian points that overlap with the rendered pixel v :

$$C(v) = \sum_{i \in \mathcal{N}} c_i \alpha_i \prod_{j=1}^{i-1} (1 - \alpha_j), \quad (3)$$

where $\mathcal{N}$ represents the ordered Gaussian points within a tile, and $\alpha_i = o_i G_i^{2D}(v)$. The $G_i^{2D}(\cdot)$ means the projected 2D Gaussian distribution of the i -th 3D Gaussian point.

As 3D Gaussian Splatting exhibits excellent real-time rendering speed while maintaining high rendering quality even at 1080p resolution, many methods [31, 30, 63] have adopted it as a 3D scene modeling technology to accelerate the rendering speed. LangSplat aims to build a 3D language field to support open-vocabulary querying within 3D spaces. It extends each 3D Gaussian with a language embedding and supervises the 3D language Gaussians with 2D CLIP image features. LangSplat presents two main contributions to make it faster and more accurate. First, instead of using multiple patch-wise CLIP features with varying scales, which leads to imprecise and vague 3D language fields, LangSplat employs the CLIP features of SAM masks with three pre-defined SAM hierarchical semantic scales to obtain precise language fields with clear boundaries. Second, as CLIP embeddings are high-dimensional, directly splatting at CLIP feature space poses huge challenges in training memory and time cost. Therefore, LangSplat proposes a scene-specific autoencoder, which compresses the high-dimensional (512-D) CLIP features into low-dimensional (3-D) latent features. The rendering process of LangSplat follows:

$$\mathbf{F}(v) = \sum_{i \in \mathcal{N}} \mathbf{f}_i \alpha_i \prod_{j=1}^{i-1} (1 - \alpha_j), \quad (4)$$

where $\mathbf{f}_i \in \mathbb{R}^d$ represents the augmented d -dimensional latent feature at i -th 3D Gaussian point and $\mathbf{F}(v)$ denotes the rendered latent feature at pixel v . After obtaining the rendered feature, LangSplat uses the decoder $g_d(\mathbf{F}(v)) \in \mathbb{R}^D$ from the trained autoencoder g to decode the latent feature back to CLIP feature space. In the end, the LangSplat attains an accurate 3D language field while being $199 \times$ faster than the previous state-of-the-art method LERF [13].

3.2 Bottleneck Analysis

While LangSplat achieved significant speedup over other methods, it still cannot perform real-time open-vocabulary 3D querying at high resolution. However, there is a high demand for real-time 3D querying for many applications, such as Augmented Reality (AR) [64], intelligent robotics [13].

To further improve the query speed of LangSplat, we first analyze each step of the query process and identify the bottleneck. For a given text query, the inference of LangSplat can be divided into three stages: rendering, decoding, and post-processing. The rendering stage will render the learned 3D language Gaussians into a 2D image plane and get a rendered latent feature map $\mathbf{F} \in \mathbb{R}^{H \times W \times d}$, where H, W denote the height and width of the 2D image size, respectively. As the rendered feature map only encodes language features in the low-dimensional latent space, a decoding stage is followed to obtain the feature map at CLIP space with a decoder $g_d(\mathbf{F}) \in \mathbb{R}^{H \times W \times D}$, where the decoder $g_d(\cdot)$ is implemented with an MLP. The last post-processing stage computes the relevancy score with the obtained D -dimensional feature following LERF [13] and remove some noise in the relevancy score map. Specifically, a mean filter is applied to the relevancy score map. Note that LangSplat adopts the three semantic levels introduced by SAM [65], so the above operations are repeated three times for three SAM semantic levels and the post-processing stage needs to select one semantic level for prediction with some specific strategies [12].

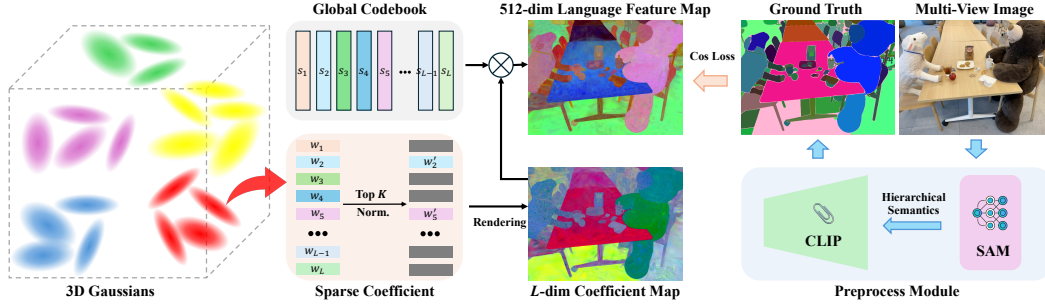


Figure 2: The framework of LangSplatV2. LangSplatV2 introduces a sparse coefficient for each Gaussian point and a shared global codebook for the entire scene.

We test the stage-wise inference time of LangSplat on the LERF dataset with one A100 GPU. The results in Table 1 show that the rendering stage takes 6.0 ms, the decoding stage takes 83.1 ms, and the post-processing stage takes 33.0 ms. While all stages occupy a considerable amount of time, we can significantly improve them through two simple modifications. First, we notice that LangSplat implements the post-processing stage on the CPU, which could be slow when it involves mean filter operations. We implement the post-processing stage on the GPU and observe that the time cost of the post-processing stage is now negligible compared with other stages. The second modification is scale parallelization. LangSplat performs text querying on three semantic levels sequentially, meaning the rendering stage is performed three times separately. However, we can perform inference at three semantic scales in parallel. For example, instead of splatting d -dimensional features three times, we directly splat $3d$ -dimensional features for once. We use LangSplat* to denote the version with these simple modifications. We observe that for LangSplat*, the rendering stage takes 2.0 ms, and the post-processing stage takes 0.5 ms. In these three stages, the decoding phase occupies 97.1% of the total query time, making it the primary bottleneck in achieving real-time 3D querying.

3.3 3D Sparse Coefficient Field

Different from 3D Gaussian Splatting, which renders 3-dimensional RGB color, 3D language fields need to splat high-dimensional features. Directly splatting high-dimensional features significantly decreases the rendering speed, as shown in Figure 1. To address this issue, existing methods usually model d -dimensional ($d \ll D$) latent features in 3D followed by either an online decoder [51] or an offline decoder [12] to decode latent features back to D -dimensional features in 2D image space. The high dimension gap between latent features and decoded features implies a heavyweight MLP is required to ensure the accuracy of the decoding stage, which becomes the primary speed bottleneck.

In LangSplatV2, instead of splatting the d -dimensional latent language features, we propose to model a 3D sparse coefficient field. As shown in Eq. 4, LangSplat assigns each Gaussian point with a language feature f_i , which represents the semantics associated with the Gaussian point. As LangSplat could create millions of Gaussian points, there will be millions of unique language features. However, this is highly inefficient as the unique semantics within a scene are quite limited and much smaller than the number of Gaussian points. As a matter of fact, many Gaussian points share similar semantics. Therefore, we assume the language embedding of every Gaussian point within a scene can be represented as a sparse coding of L global basis vectors $\mathcal{S} = [s_1, s_2, \dots, s_L]^T \in \mathbb{R}^{L \times D}$. These L basis vectors serve as the global codebook and each Gaussian point is computed by linearly combining a small number of local basis vectors. We assume that only K ($K \ll L$) basis vectors from L global codebook are used to represent the language embedding of a Gaussian point. Then we define $w_i = [w_{i,1}, w_{i,2}, \dots, w_{i,L}]^T \in \mathbb{R}^{1 \times L}$ as the associated sparse coefficients for i -th Gaussian point, where $\sum_{l=1}^L w_{i,l} = 1$ and only K elements are non-zero values while all other elements are zeros. With slightly abuse of the notion f_i , the language embedding $f_i \in \mathbb{R}^D$ of i -th Gaussian point is represented as:

$$f_i = w_i \mathcal{S} = \sum_{l=1}^L w_{i,l} s_l. \quad (5)$$

If we directly render the D -dimensional language field without compressing CLIP features, we have:

$$\mathbf{S} = \sum_{i \in \mathcal{N}} \mathbf{w}_i \mathcal{S} \alpha_i \prod_{j=1}^{i-1} (1 - \alpha_j), \quad (6)$$

where $\mathbf{S}$ is the rendered D -dimensional CLIP feature. We set $e_i = \alpha_i \prod_{j=1}^{i-1} (1 - \alpha_j)$, then we have:

$$\mathbf{S} = \sum_{i \in \mathcal{N}} \mathbf{w}_i \mathcal{S} e_i = \left(\sum_{i \in \mathcal{N}} e_i \mathbf{w}_i \right) \mathcal{S}. \quad (7)$$

Eq. 7 shows that rendering the D -dimensional CLIP features is equivalent to first rendering the sparse coefficients $\mathbf{w}(i)$, and then performing a matrix multiplication with the global dictionary $\mathcal{S}$.

Therefore, we propose to learn a L -dimensional sparse coefficient for each Gaussian point and a global codebook with a size of L basis vectors. To learn a sparse distribution for the coefficient $\mathbf{w}_i$, we first apply a softmax function to normalize the L -dimensional parameter, then preserve the top- K values as non-zeros elements and set the remaining elements to zeros. We will re-normalize the top- K values with their sum to ensure the sum of the coefficient equals one. The L -dimensional 3D sparse coefficient field and the L basis vectors are jointly learned. Figure 2 visualizes the framework.

Compared to LangSplat, LangSplatV2 requires only a simple matrix multiplication after rendering the weight map, instead of relying on a heavyweight decoder, thus overcoming LangSplat’s main inference speed bottleneck. Furthermore, the D -dimensional global codebook eliminates reconstruction loss from dimensionality reduction, enabling better modeling of high-dimensional features in the CLIP space and improving query accuracy.

3.4 Efficient Sparse Coefficient Splatting

By learning a 3D sparse coefficient field, the MLP decoder is entirely removed, eliminating the associated computational overhead of the MLP. However, we still need to perform a $3L$ -dimensional feature rendering and a matrix multiplication. Our experiments show that the time for matrix multiplication (in Table 1, listed as the decoding stage of LangSplatV2) is negligible compared to the rendering process. However, rendering high-dimensional features with dimensionality L remains computationally demanding. To address this issue, we propose an efficient sparse coefficient splatting method with CUDA optimization. By exploiting the sparsity of the coefficient field, we can achieve L -dimensional feature rendering at the cost of only K -dimensional rendering, where $K \ll L$.

In the CUDA implementation of 3D Gaussian Splatting and LangSplat, each thread performs alpha-blending of the $|\mathcal{N}|$ ordered Gaussian points within a tile. With an L -dimensional rendering, each thread sequentially computes L channels, leading to a computational complexity of $O(|\mathcal{N}|L)$. When L becomes sufficiently large, this alpha-blending computation becomes the key bottleneck and significantly increases the overall computational overhead as L grows, as shown in Figure 1.

To mitigate this, we utilize the sparse nature of the learned coefficient field during test-time querying, as shown in Figure 3. Although each Gaussian point has an L -dimensional coefficient, only K dimensions contain non-zero values. Therefore we can only perform alpha-blending for the non-zero elements, as the blending of zero elements does not alter the result. This effectively reduces the computational complexity from $O(|\mathcal{N}|L)$ to $O(|\mathcal{N}|K)$, with K significantly smaller than L . In practice, we set $K = 4$, which allows us to simultaneously render three semantic scales with an effective feature rendering dimension in CUDA of only 12, yielding a high-quality feature map of 1,536 dimensions. It makes the rendering highly efficient without compromising feature quality.

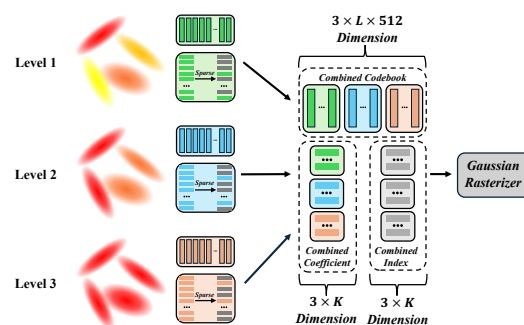


Figure 3: Our efficient sparse coefficient splatting method accelerates the speed of alpha-blending by utilizing the property of the learned sparse coefficient field and neglecting zero elements.

Table 2: Quantitative comparisons of open-vocabulary 3D object localization and 3D semantic segmentation on the LERF dataset. We report the mean accuracy (%) and the mean IoU scores (%).

Method	3D Object Localization					3D Semantic Segmentation				
	Ramen	Teatime	Kitchen	Figurines	Overall	Ramen	Teatime	Kitchen	Figurines	Overall
GS-Grouping [67]	32.4	69.5	50.0	44.6	49.1	26.4	54.0	31.3	34.6	36.6
LEGaussian [49]	69.0	79.7	63.6	57.1	67.4	20.2	32.3	22.3	23.4	24.6
GOI [52]	56.3	67.8	68.2	44.6	59.2	33.7	55.8	54.5	23.9	42.0
GAGS [15]	69.0	<u>88.1</u>	<u>90.9</u>	78.6	81.7	46.8	60.3	<u>55.8</u>	<u>53.6</u>	<u>54.1</u>
LangSplat [12]	<u>73.2</u>	<u>88.1</u>	95.5	<u>80.4</u>	84.3	<u>51.2</u>	<u>65.1</u>	44.5	44.7	51.4
LangSplatV2	74.7	93.2	86.4	82.1	<u>84.1</u>	51.8	72.2	59.1	56.4	59.9

Specifically, each Gaussian’s L -dimensional sparse coefficient w_i can be fully represented by its top- K non-zero elements. We store these as two K -dimensional arrays for each Gaussian: top- K indices and top- K coefficients. The top- K indices are the positions of the top- K non-zero elements within the L -dimensional vector and the top- K coefficients are the values of these top- K non-zero elements. During rendering, each CUDA thread performs alpha-blending only for the K non-zero coefficients. For each Gaussian point within the tile, the CUDA thread will access the indices and coefficients of the top- K elements and perform weighted summation solely on these K -dimensional indices, avoiding computation on zero elements. In this way, LangSplatV2 can obtain high-dimensional (D) feature splatting results at the cost of only ultra-low-dimensional (K) feature splatting.

4 Experiments

4.1 Datasets and Details

Datasets. We evaluate our method on the LERF, 3D-OVS, and Mip-NeRF360 datasets. The LERF dataset [13], captured using the iPhone App Polycam, contains in-the-wild scenes. For the open-vocabulary 3D object localization task, we adopt the augmented localization annotations provided by LangSplat [12] on the LERF dataset. Additionally, we use the segmentation ground truth from LangSplat [12] for the open-vocabulary 3D segmentation task on LERF. Beyond LERF, we also conduct 3D segmentation experiments on the 3D-OVS and Mip-NeRF360 [66] datasets. The 3D-OVS dataset [47] includes a collection of long-tail objects captured in diverse poses and backgrounds. Moreover, we evaluate our method on Mip-NeRF360, which consists of multi-view indoor and outdoor scene images, with segmentation labels annotated by GAGS [15]. For evaluation, we use localization accuracy for the 3D object localization task and report the average IoU scores for the open-vocabulary 3D segmentation task.

Implementation Details. Following LangSplat [12], we use the OpenCLIP ViT-B/16 model to extract CLIP features. We employ the ViT-H model for SAM [65] to segment images and obtain masks with three hierarchical semantics. The codebook size L is set to 64 and the K is set to 4. During test-time querying, we render three semantic scales simultaneously, leading to the actual rendering dimension of 12. The 3D Gaussians are first trained with RGB supervision for 30,000 iterations to reconstruct the RGB scene. Then we train another 10,000 iterations for the 3D sparse coefficient field by fixing all other 3D Gaussian parameters. All our experiments are conducted on one A100 GPU.

4.2 Quantitative Results

Time Analysis. To assess the speed improvement of LangSplatV2 over its LangSplat, we conducted a detailed time analysis using the LERF dataset, with the computations performed on a single A100 GPU. Table 1 presents a stage-wise breakdown of the time costs for LangSplat, LangSplat*, and LangSplatV2. LangSplat, our initial model, achieves a frame rate of 8.2 FPS. The LangSplat* improvements mainly focused on rendering and post-processing optimizations, which reduced the total time to 85.6 ms and increased the speed to 11.7 FPS. Here, the rendering time was significantly cut to 2.0 ms, and the post-processing time was reduced to less than 1.0 ms with our proposed simple engineering modifications. In contrast, LangSplatV2 proposed an efficient sparse coefficient splatting method, that entirely removed the MLP decoder in the decoding stage. The only operation in the

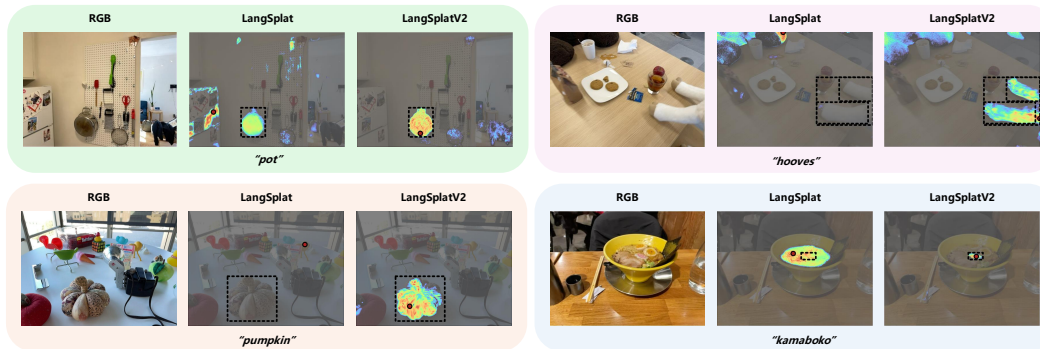


Figure 4: Qualitative comparisons of open-vocabulary 3D object localization on the LERF dataset. The red points are the model predictions and the black dashed bounding boxes denote the annotations. We observe that LangSplatV2 generates better results than LangSplat.

decoding stage is performing a matrix multiplication between the rendered L -dimensional coefficient and the global dictionary, which takes only 0.1 ms. While completely removing the MLP decoder, our LangSplatV2 only needs 2.0 ms to render the sparse coefficient, leading to an overall 476.2 FPS for obtaining high-dimensional feature maps. In the end, LangSplatV2 achieved an exceptional reduction in total querying time to 2.6 ms. These results not only demonstrate a substantial enhancement in processing speed but also affirm the real-time capabilities of LangSplatV2 for high-resolution, complex 3D scene querying.

Main Results on the LERF dataset.

Table 2 demonstrates the comparisons on the LERF dataset. For the open-vocabulary 3D object localization task, LangSplatV2 achieved the highest accuracy with notable improvements over previous methods for most scenes, including “ramen”, “figurines”, and “teatime”. In the “Kitchen” scene, LangSplat achieved a higher accuracy than LangSplatV2. This difference is mainly due to the limited sample size (22 examples in total), with only a two-prediction gap (21 vs. 19). Given LangSplat’s high variance in this scene, we ran it multiple times, obtaining an average of 17.5 ± 1.8 correct predictions, which suggests that LangSplatV2 provides more consistent performance advantage across different scenarios. Results on 3D semantic segmentation demonstrate that LangSplatV2 outperforms all other methods, and it consistently surpasses LangSplat across all scenes, highlighting the effectiveness of our proposed method.

Main Results on the 3D-OVS dataset.

Table 3 shows the quantitative comparisons of open-vocabulary 3D semantic segmentation performance, on the 3D-OVS dataset across five test scenes. We see that LangSplatV2 achieves the highest overall mean IoU score of 94.6%. It significantly outperforms LangSplat across all different scenes, further validating its ability to build more accurate language fields.

Main Results on the Mip-NeRF360 dataset. Table 4 presents the comparison results on the Mip-NeRF360 dataset. Notably, LangSplatV2 achieved a significant improvement in the overall average IoU score, reaching 69.4%, which is a marked advancement over LangSplat’s 57.3%.

Table 3: Quantitative 3D semantic segmentation results on 3D-OVS, reported as mean IoU (%).

Method	Bed	Bench	Room	Sofa	Lawn	Overall
FFD [1]	56.6	6.1	25.1	3.7	42.9	26.9
LERF [13]	73.5	53.2	46.6	27.0	73.7	54.8
3D-OVS [47]	89.5	89.3	92.8	74.0	88.2	86.8
GS-Grouping [67]	83.0	91.5	85.9	87.3	90.6	87.7
LEGaussian [49]	84.9	91.1	86.0	87.8	92.5	88.5
GOI [52]	89.4	92.8	91.3	85.6	94.1	90.6
LangSplat [12]	<u>92.5</u>	<u>94.2</u>	<u>94.1</u>	<u>90.0</u>	<u>96.1</u>	<u>93.4</u>
LangSplatV2	93.0	94.9	96.1	92.3	96.6	94.6

Table 4: Quantitative 3D semantic segmentation results on Mip-NeRF360, reported as mean IoU (%).

Method	Room	Counter	Garden	Bonsai	Overall
GS-Grouping [67]	54.4	47.7	40.4	54.1	49.2
LEGaussian [49]	25.5	35.3	33.2	22.3	29.1
GOI [52]	60.3	46.6	59.8	67.3	58.5
GAGS [15]	65.2	61.1	<u>61.2</u>	<u>70.5</u>	<u>64.5</u>
LangSplat [12]	53.2	<u>68.8</u>	51.9	55.4	57.3
LangSplatV2	<u>64.3</u>	75.1	65.0	73.1	69.4

Table 5: Ablation study on codebook size L . We report the average performance of the 3D object localization and 3D semantic segmentation tasks on the LERF dataset.

L	32	64	128
Accuracy (%)	72.8	84.1	84.1
IoU (%)	53.9	59.9	60.5

Table 6: Ablation study on different K . We report the average performance of the 3D object localization and 3D semantic segmentation tasks on the LERF dataset.

K	2	4	6	8
Accuracy (%)	79.4	84.1	84.4	83.6
IoU (%)	54.4	59.9	59.9	59.9

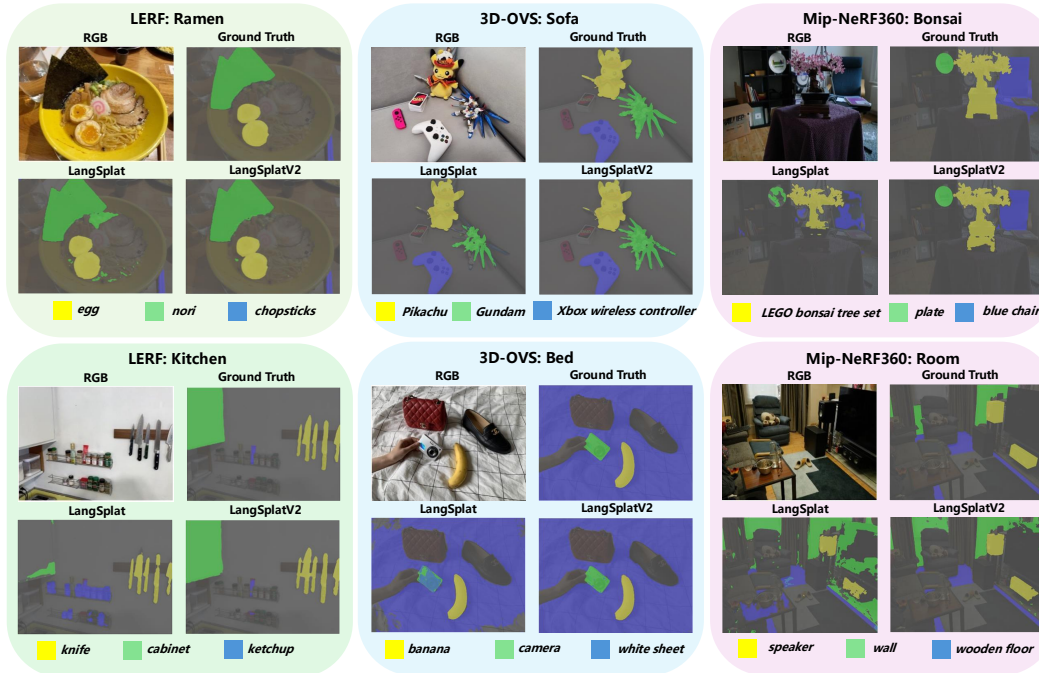


Figure 5: Qualitative comparisons of open-vocabulary 3D semantic segmentation on the LERF, Mip-NeRF360 and 3D-OVS dataset. We can see that our LangSplatV2 generates better masks than LangSplat, which shows the effectiveness of our LangSplatV2.

Ablation Study. In Table 5, we evaluate the effect of codebook sizes L on the localization and segmentation tasks. We report the mean accuracy and IoU scores on the LERF dataset. Our results reveal that increasing L from 32 to 64 yields a notable improvement in both localization accuracy (72.8% to 84.1%) and segmentation IoU (53.9% to 59.9%). This suggests that a larger codebook size better captures the semantic fields in complex scenes, allowing for more precise language representations. However, further increasing L results in only slightly increased performance, suggesting that $L = 64$ has already reached saturation on the LERF dataset. Table 6 further reports the effects of different K . We observe that $K = 4$ can saturate the performance of the Gaussian model, and a larger K will increase the rendering dimensions, thereby slowing down the rendering speed. More ablation studies and details can be found in the supplementary material.

Training Cost. Table 7 shows the comparison of training cost on the LERF [13] dataset with one A100 GPU. Compared with LEGaussians [49] and LangSplat [12], LangSplatV2 comes with increased training cost due to the need to construct high-dimensional semantic fields during training. Although our LangSplatV2 incurs higher training costs compared to LangSplat and LEGaussians, our primary focus in this paper is on improving the model’s test-time performance for deployment. As demonstrated in the experiments, LangSplatV2 achieves significantly higher inference speed— $47\times$ faster than LangSplat and $14\times$ faster than LEGaussians on the LERF dataset—under comparable memory consumption, while also yielding higher segmentation accuracy. Furthermore, compared with naive LangSplat (directly trained with 512 dimensions), LangSplatV2 is more efficient with the proposed global codebook and sparse coefficient field.

Discussion. Recent work has also explored codebook-based language field representations, such as LEGaussians [49]. However, our approach departs significantly from theirs in several key aspects.

Table 7: Quantitative comparisons of training cost on the LERF dataset.

Method	LangSplat [12]	LEGaussian [49]	LangSplatV2	LangSplat-512D [12]
Training Time (h)	1.0	1.3	3.0	45.0
Training Memory (GB)	6.2	11	21.2	47.4

Table 8: Quantitative comparisons with LEGaussian on three benchmarks.

Method	Segmentation IoU (%)			Query Time (ms)			GPU Memory (GB)		
	LERF	3DOVS	Mip.	LERF	3DOVS	Mip.	LERF	3DOVS	Mip.
LEGaussian [49]	24.6	88.5	29.1	36.7	59.6	58.0	8.2	10.5	14.0
LangSplatV2	59.9	94.6	69.4	2.6	4.8	3.8	7.2	11.6	16.9

First, LEGaussians construct a 2D feature codebook from 2D images while LangSplatV2 directly learns a global codebook in 3D space, which is more efficient. Second, LEGaussians still rely on an MLP to compress features and do not exploit sparsity. In contrast, our method completely removes the MLP and further enforces and leverages sparsity in both representation and rendering. As shown in Table 8, our method achieves superior semantic segmentation accuracy, faster inference speed, and comparable memory usage across three benchmark datasets, which demonstrates the effectiveness of our approach. For a more detailed comparison and discussion, please refer to the Appendix.

4.3 Qualitative Results

Figure 4 visualizes the open-vocabulary 3D object localization results. We observe that our LangSplatV2 can give more accurate localization predictions compared with LangSplat. For example, our LangSplatV2 can give the correct prediction for “pumpkin” while LangSplat entirely fails in this query. Figure 5 visualizes the open-vocabulary 3D segmentation results in three datasets. We observe that our LangSplatV2 can generate more accurate masks compared with LangSplat. For example, in the “Kitchen” scene, LangSplat predicts a noisy mask for the “ketchup” query, while our LangSplatV2 generates more precise and clean masks.

5 Conclusion

Conclusion. In this paper, we have presented LangSplatV2 for high-dimensional language feature splatting in 3D open-vocabulary querying, designed to overcome the speed limitations of LangSplat. Our analysis identified the decoding stage as the primary bottleneck in LangSplat. By modeling each Gaussian point as a sparse code over a global dictionary, LangSplatV2 removes the need for a heavyweight decoder. We further proposed an efficient sparse coefficient splatting method with CUDA optimization, allowing LangSplatV2 to achieve high-dimensional feature splatting results at the cost of splatting ultra-low-dimensional features.

Limitations and broader impacts. While LangSplatV2 demonstrates improved performance and faster inference compared to LangSplat, it comes with increased training cost due to the need to construct high-dimensional semantic fields during training. Furthermore, as our method directly inherits the semantic representations from the CLIP model, it also inherits its inherent biases. Addressing such biases remains an open research challenge and may benefit from recent advances in fairness-aware versions of CLIP.

Acknowledgment

This research is supported in part by the NIH grant R01HD104969.

References

- [1] Sosuke Kobayashi, Eiichi Matsumoto, and Vincent Sitzmann. Decomposing nerf for editing via feature field distillation. *Advances in Neural Information Processing Systems*, 35:23311–23330, 2022.

- [2] Vadim Tschernezki, Iro Laina, Diane Larlus, and Andrea Vedaldi. Neural feature fusion fields: 3d distillation of self-supervised 2d image representations. In *2022 International Conference on 3D Vision (3DV)*, pages 443–453. IEEE, 2022.
- [3] Alan B Craig. Understanding augmented reality: Concepts and applications. 2013.
- [4] Ronald T Azuma. A survey of augmented reality. *Presence: teleoperators & virtual environments*, 6(4):355–385, 1997.
- [5] Songyou Peng, Kyle Genova, Chiyu Jiang, Andrea Tagliasacchi, Marc Pollefeys, Thomas Funkhouser, et al. Openscene: 3d scene understanding with open vocabularies. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 815–824, 2023.
- [6] Chenguang Huang, Oier Mees, Andy Zeng, and Wolfram Burgard. Visual language maps for robot navigation. In *2023 IEEE International Conference on Robotics and Automation (ICRA)*, pages 10608–10615. IEEE, 2023.
- [7] Paola Cascante-Bonilla, Hui Wu, Letao Wang, Rogerio S Feris, and Vicente Ordonez. Simvqa: Exploring simulated environments for visual question answering. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 5056–5066, 2022.
- [8] Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. Learning transferable visual models from natural language supervision. In *International conference on machine learning*, pages 8748–8763. PMLR, 2021.
- [9] Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae Lee. Visual instruction tuning. *Advances in neural information processing systems*, 36, 2024.
- [10] Ben Mildenhall, Pratul P Srinivasan, Matthew Tancik, Jonathan T Barron, Ravi Ramamoorthi, and Ren Ng. Nerf: Representing scenes as neural radiance fields for view synthesis. *Communications of the ACM*, 65(1):99–106, 2021.
- [11] Bernhard Kerbl, Georgios Kopanas, Thomas Leimkühler, and George Drettakis. 3d gaussian splatting for real-time radiance field rendering. *ACM Transactions on Graphics*, 42(4):1–14, 2023.
- [12] Minghan Qin, Wanhua Li, Jiawei Zhou, Haoqian Wang, and Hanspeter Pfister. Langsplat: 3d language gaussian splatting. In *CVPR*, 2024.
- [13] Justin Kerr, Chung Min Kim, Ken Goldberg, Angjoo Kanazawa, and Matthew Tancik. Lerf: Language embedded radiance fields. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 19729–19739, 2023.
- [14] Ola Shorinwa, Jiankai Sun, and Mac Schwager. Fast-splat: Fast, ambiguity-free semantics transfer in gaussian splatting. *arXiv preprint arXiv:2411.13753*, 2024.
- [15] Yuning Peng, Haiping Wang, Yuan Liu, Chenglu Wen, Zhen Dong, and Bisheng Yang. Gags: Granularity-aware 3d feature distillation for gaussian splatting. *arXiv preprint arXiv:2412.13654*, 2024.
- [16] Raymond Koon Chuan Koh, Henry Been-Lirn Duh, and Jian Gu. An integrated design flow in user interface and interaction for enhancing mobile ar gaming experiences. In *2010 IEEE International Symposium on Mixed and Augmented Reality-Arts, Media, and Humanities*, pages 47–52. IEEE, 2010.
- [17] Robert Godwin-Jones. Augmented reality and language learning: From annotated vocabulary to place-based mobile games. 2016.
- [18] Zehao Yu, Anpei Chen, Binbin Huang, Torsten Sattler, and Andreas Geiger. Mip-splatting: Alias-free 3d gaussian splatting. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 19447–19456, 2024.
- [19] Yang Fu, Sifei Liu, Amey Kulkarni, Jan Kautz, Alexei A Efros, and Xiaolong Wang. Colmap-free 3d gaussian splatting. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 20796–20805, 2024.
- [20] Jixuan Fan, Wanhua Li, Yifei Han, Tianru Dai, and Yansong Tang. Momentum-gs: Momentum gaussian self-distillation for high-quality large scene reconstruction. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 25250–25260, 2025.

- [21] Haonan Han, Rui Yang, Huan Liao, Jiankai Xing, Zunnan Xu, Xiaoming Yu, Junwei Zha, Xiu Li, and Wanhua Li. Reparo: Compositional 3d assets generation with differentiable 3d layout alignment. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 25367–25377, 2025.
- [22] Heng Yu, Joel Julin, Zoltán Á Milacski, Koichiro Niinuma, and László A Jeni. Cogs: Controllable gaussian splatting. *arXiv preprint arXiv:2312.05664*, 2023.
- [23] Xian Liu, Xiaohang Zhan, Jiayang Tang, Ying Shan, Gang Zeng, Dahua Lin, Xihui Liu, and Ziwei Liu. Humangaussian: Text-driven 3d human generation with gaussian splatting. *arXiv preprint arXiv:2311.17061*, 2023.
- [24] Zhicheng Lu, Xiang Guo, Le Hui, Tianrui Chen, Min Yang, Xiao Tang, Feng Zhu, and Yuchao Dai. 3d geometry-aware deformable gaussian splatting for dynamic view synthesis. *arXiv preprint arXiv:2404.06270*, 2024.
- [25] Kangjie Chen, Yingji Zhong, Zhihao Li, Jiaqi Lin, Youyu Chen, Minghan Qin, and Haoqian Wang. Quantifying and alleviating co-adaptation in sparse-view 3d gaussian splatting. *arXiv preprint arXiv:2508.12720*, 2025.
- [26] Yuanhao Cai, Zihao Xiao, Yixun Liang, Minghan Qin, Yulun Zhang, Xiaokang Yang, Yaoyao Liu, and Alan L Yuille. Hdr-gs: Efficient high dynamic range novel view synthesis at 1000x speed via gaussian splatting. *Advances in Neural Information Processing Systems*, 37:68453–68471, 2024.
- [27] Dongbin Zhang, Yunfei Liu, Lijian Lin, Ye Zhu, Yang Li, Minghan Qin, Yu Li, and Haoqian Wang. Guava: Generalizable upper body 3d gaussian avatar. *arXiv preprint arXiv:2505.03351*, 2025.
- [28] Yang Liu, Xiang Huang, Minghan Qin, Qinwei Lin, and Haoqian Wang. Animatable 3d gaussian: Fast and high-quality reconstruction of multiple human avatars. In *Proceedings of the 32nd ACM International Conference on Multimedia*, pages 1120–1129, 2024.
- [29] Dongbin Zhang, Chuming Wang, Weitao Wang, Peihao Li, Minghan Qin, and Haoqian Wang. Gaussian in the wild: 3d gaussian splatting for unconstrained image collections. In *European Conference on Computer Vision*, pages 341–359. Springer, 2024.
- [30] Ziyi Yang, Xinyu Gao, Wen Zhou, Shaohui Jiao, Yuqing Zhang, and Xiaogang Jin. Deformable 3d gaussians for high-fidelity monocular dynamic scene reconstruction. *arXiv preprint arXiv:2309.13101*, 2023.
- [31] Guanjun Wu, Taoran Yi, Jiemin Fang, Lingxi Xie, Xiaopeng Zhang, Wei Wei, Wenyu Liu, Qi Tian, and Xinggang Wang. 4d gaussian splatting for real-time dynamic scene rendering. In *CVPR*, 2024.
- [32] Yiwen Chen, Zilong Chen, Chi Zhang, Feng Wang, Xiaofeng Yang, Yikai Wang, Zhongang Cai, Lei Yang, Huaping Liu, and Guosheng Lin. Gaussianeditor: Swift and controllable 3d editing with gaussian splatting. *arXiv preprint arXiv:2311.14521*, 2023.
- [33] Jiemin Fang, Junjie Wang, Xiaopeng Zhang, Lingxi Xie, and Qi Tian. Gaussianeditor: Editing 3d gaussians delicately with text instructions. *arXiv preprint arXiv:2311.16037*, 2023.
- [34] Ayaan Haque, Matthew Tancik, Alexei A Efros, Aleksander Holynski, and Angjoo Kanazawa. Instruct-nerf2nerf: Editing 3d scenes with instructions. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 19740–19750, 2023.
- [35] Shijie Zhou, Zhiwen Fan, Dejia Xu, Haoran Chang, Pradyumna Chari, Tejas Bharadwaj, Suyu You, Zhangyang Wang, and Achuta Kadambi. Dreamscene360: Unconstrained text-to-3d scene generation with panoramic gaussian splatting. *arXiv preprint arXiv:2404.06903*, 2024.
- [36] Jiayang Tang, Zhaoxi Chen, Xiaokang Chen, Tengfei Wang, Gang Zeng, and Ziwei Liu. Lgm: Large multi-view gaussian model for high-resolution 3d content creation. *arXiv preprint arXiv:2402.05054*, 2024.
- [37] Jiawei Ren, Liang Pan, Jiayang Tang, Chi Zhang, Ang Cao, Gang Zeng, and Ziwei Liu. Dreamgaussian4d: Generative 4d gaussian splatting. *arXiv preprint arXiv:2312.17142*, 2023.
- [38] Zilong Chen, Feng Wang, and Huaping Liu. Text-to-3d using gaussian splatting. *arXiv preprint arXiv:2309.16585*, 2023.
- [39] Fangfu Liu, Wenqiang Sun, Hanyang Wang, Yikai Wang, Haowen Sun, Junliang Ye, Jun Zhang, and Yueqi Duan. Reconx: Reconstruct any scene from sparse views with video diffusion model. *arXiv preprint arXiv:2408.16767*, 2024.

- [40] Hanyang Wang, Fangfu Liu, Jiawei Chi, and Yueqi Duan. Videoscene: Distilling video diffusion model to generate 3d scenes in one step. In *2025 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 16475–16485. IEEE, 2025.
- [41] William Shen, Ge Yang, Alan Yu, Jansen Wong, Leslie Pack Kaelbling, and Phillip Isola. Distilled feature fields enable few-shot language-guided manipulation. *arXiv preprint arXiv:2308.07931*, 2023.
- [42] Yawar Siddiqui, Lorenzo Porzi, Samuel Rota Buló, Norman Müller, Matthias Nießner, Angela Dai, and Peter Kotschieder. Panoptic lifting for 3d scene understanding with neural fields. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 9043–9052, 2023.
- [43] Shuaifeng Zhi, Tristan Laidlow, Stefan Leutenegger, and Andrew J Davison. In-place scene labelling and understanding with implicit scene representation. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 15838–15847, 2021.
- [44] Hao Li, Roy Qin, Zhengyu Zou, Diqi He, Bohan Li, Bingquan Dai, Dingwen Zhang, and Junwei Han. Langsurf: Language-embedded surface gaussians for 3d scene understanding. *arXiv preprint arXiv:2412.17635*, 2024.
- [45] Kangjie Chen, BingQuan Dai, Minghan Qin, Dongbin Zhang, Peihao Li, Yingshuang Zou, and Haoqian Wang. Slgaussian: Fast language gaussian splatting in sparse views. *arXiv preprint arXiv:2412.08331*, 2024.
- [46] Wanhua Li, Renping Zhou, Jiawei Zhou, Yingwei Song, Johannes Herter, Minghan Qin, Gao Huang, and Hanspeter Pfister. 4d langplat: 4d language gaussian splatting via multimodal large language models. In *Proceedings of the Computer Vision and Pattern Recognition Conference*, pages 22001–22011, 2025.
- [47] Kunhao Liu, Fangneng Zhan, Jiahui Zhang, Muyu Xu, Yingchen Yu, Abdulmotaleb El Saddik, Christian Theobalt, Eric Xing, and Shijian Lu. Weakly supervised 3d open-vocabulary segmentation. *Advances in Neural Information Processing Systems*, 36:53433–53456, 2023.
- [48] Mathilde Caron, Hugo Touvron, Ishan Misra, Hervé Jégou, Julien Mairal, Piotr Bojanowski, and Armand Joulin. Emerging properties in self-supervised vision transformers. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 9650–9660, 2021.
- [49] Jin-Chuan Shi, Miao Wang, Hao-Bin Duan, and Shao-Hua Guan. Language embedded 3d gaussians for open-vocabulary scene understanding. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 5333–5343, 2024.
- [50] Xingxing Zuo, Pouya Samangouei, Yunwen Zhou, Yan Di, and Mingyang Li. Fmgs: Foundation model embedded 3d gaussian splatting for holistic 3d scene understanding. *arXiv preprint arXiv:2401.01970*, 2024.
- [51] Shijie Zhou, Haoran Chang, Sicheng Jiang, Zhiwen Fan, Zehao Zhu, Dejia Xu, Pradyumna Chari, Suyu You, Zhangyang Wang, and Achuta Kadambi. Feature 3dgs: Supercharging 3d gaussian splatting to enable distilled feature fields. *arXiv preprint arXiv:2312.03203*, 2023.
- [52] Yansong Qu, Shaohui Dai, Xinyang Li, Jianghang Lin, Liujuan Cao, Shengchuan Zhang, and Rongrong Ji. Goi: Find 3d gaussians of interest with an optimizable open-vocabulary semantic-space hyperplane, 2024.
- [53] Simon Niedermayr, Josef Stumpfegger, and Rüdiger Westermann. Compressed 3d gaussian splatting for accelerated novel view synthesis. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 10349–10358, 2024.
- [54] KL Navaneet, Kossar Pourahmadi Meibodi, Soroush Abbasi Koohpayegani, and Hamed Pirsiavash. Compgs: Smaller and faster gaussian splatting with vector quantization. In *European Conference on Computer Vision*, pages 330–349. Springer, 2024.
- [55] Wieland Morgenstern, Florian Barthel, Anna Hilsmann, and Peter Eisert. Compact 3d scene representation via self-organizing gaussian grids. In *European Conference on Computer Vision*, pages 18–34. Springer, 2024.
- [56] Zhiwen Fan, Kevin Wang, Kairun Wen, Zehao Zhu, Dejia Xu, Zhangyang Wang, et al. Lightgaussian: Unbounded 3d gaussian compression with 15x reduction and 200+ fps. *Advances in neural information processing systems*, 37:140138–140158, 2024.
- [57] Joo Chan Lee, Daniel Rho, Xiangyu Sun, Jong Hwan Ko, and Eunbyung Park. Compact 3d gaussian representation for radiance field. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 21719–21728, 2024.

- [58] Alex Hanson, Allen Tu, Geng Lin, Vasu Singla, Matthias Zwicker, and Tom Goldstein. Speedy-splat: Fast 3d gaussian splatting with sparse pixels and sparse primitives. In *Proceedings of the Computer Vision and Pattern Recognition Conference*, pages 21537–21546, 2025.
- [59] Qiqi Hou, Randall Rauwendaal, Zifeng Li, Hoang Le, Farzad Farhadzadeh, Fatih Porikli, Alexei Bourd, and Amir Said. Sort-free gaussian splatting via weighted sum rendering. *arXiv preprint arXiv:2410.18931*, 2024.
- [60] Hanspeter Pfister, Matthias Zwicker, Jeroen Van Baar, and Markus Gross. Surfels: Surface elements as rendering primitives. In *Proceedings of the 27th annual conference on Computer graphics and interactive techniques*, pages 335–342, 2000.
- [61] Matthias Zwicker, Hanspeter Pfister, Jeroen Van Baar, and Markus Gross. Surface splatting. In *Proceedings of the 28th annual conference on Computer graphics and interactive techniques*, pages 371–378, 2001.
- [62] Matthias Zwicker, Hanspeter Pfister, Jeroen Van Baar, and Markus Gross. Ewa volume splatting. In *Proceedings Visualization, 2001. VIS'01.*, pages 29–538. IEEE, 2001.
- [63] Zi-Xin Zou, Zhipeng Yu, Yuan-Chen Guo, Yangguang Li, Ding Liang, Yan-Pei Cao, and Song-Hai Zhang. Triplane meets gaussian splatting: Fast and generalizable single-view 3d reconstruction with transformers. *arXiv preprint arXiv:2312.09147*, 2023.
- [64] Priyanka Jain, Nivedita Bhirud, Subhash Tatale, Abhishek Kale, Mayank Bhale, Aakanksha Hajare, and NK Jain. Real-time interactive ar for cognitive learning. In *AI, IoT, Big Data and Cloud Computing for Industry 4.0*, pages 219–239. Springer, 2023.
- [65] Alexander Kirillov, Eric Mintun, Nikhila Ravi, Hanzi Mao, Chloe Rolland, Laura Gustafson, Tete Xiao, Spencer Whitehead, Alexander C Berg, Wan-Yen Lo, et al. Segment anything. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 4015–4026, 2023.
- [66] Jonathan T. Barron, Ben Mildenhall, Dor Verbin, Pratul P. Srinivasan, and Peter Hedman. Mip-nerf 360: Unbounded anti-aliased neural radiance fields. In *2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 5460–5469, 2022.
- [67] Mingqiao Ye, Martin Danelljan, Fisher Yu, and Lei Ke. Gaussian grouping: Segment and edit anything in 3d scenes. In *European Conference on Computer Vision*, pages 162–179. Springer, 2024.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [Yes]

Justification: The main claims made in the abstract and introduction accurately reflect the paper's contributions and scope.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [Yes]

Justification: We discussed the limitations in the Conclusion Section.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [NA]

Justification: This paper does not involve theoretical contributions.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: We have included an Algorithm to demonstrate how to reproduce our method. We have released the code.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: We have released the source codes.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: We have included it in the Experiments Section.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [NA]

Justification: It is not included in all previous work in this field.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.

- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We have included it in the Experiments Section.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: Our research conducted in the paper conforms to the NeurIPS Code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: We have discussed the broader impact in the Conclusion Section.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to

generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.

- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: This paper poses no such risks.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: We have properly cited the original paper.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: This paper does not release new assets.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. Declaration of LLM usage

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The core method development in this research does not involve LLMs as any important, original, or non-standard components.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.

A Algorithm

The proposed efficient sparse coefficient splatting process is shown in Algorithm 1.

Algorithm 1 Efficient Sparse Coefficient Splatting

```
Initialize  $\mathbf{W} \in \mathbb{R}^L$  to zero
for Gaussian point  $i \in \mathcal{N}$  do
    Retrieve top- $K$  indices  $\{j_1, j_2, \dots, j_K\}$  and corresponding coefficients  $\{w_{i,j_1}, w_{i,j_2}, \dots, w_{i,j_K}\}$ 
    for each  $j \in \{j_1, \dots, j_K\}$  do
         $\mathbf{W}[j] += w_{i,j} \cdot \alpha_i \cdot \prod_{m=1}^{i-1} (1 - \alpha_m)$ 
    end for
end for
```

B Discussion

LangSplatV2 proposes to model the high-dimensional language features of each Gaussian point as a sparse code within a global dictionary. While similar concepts, such as quantizing Gaussian point attributes (*e.g.*, SH coefficients [56], opacity [53], and semantic features [49]) into a codebook, have been explored in the domains of compression [53] and language embedding [49], our method is fundamentally different.

In the compression domain, existing approaches [53, 56] first learn unique attributes for each Gaussian point, which are then quantized into a codebook using a quantizer. Extending this concept to our scenario would require learning 512-dimensional semantic features for each Gaussian point before quantization. This process significantly increases memory consumption and training time, often exceeding the capacity of GPUs like the NVIDIA 3090 or 4090, leading to out-of-memory (OOM) errors. Moreover, even with advanced GPUs that can handle the training, the rendering process would still involve managing 512-dimensional features. As shown in Figure 1, this dramatically reduces rendering speed compared to LangSplatV2.

In the context of language embedding, LEGaussians [49] adopt a codebook-based approach to accelerate training. However, our method is significantly different from LEGaussians in three key aspects: 1) Global 3D Codebook vs. 2D Codebook: LEGaussians first train a 2D codebook by quantizing features in 2D images, then learn a 3D model to predict the one-hot class category indicating the index of the codebook. In contrast, LangSplatV2 learns a global 3D codebook shared among all 3D Gaussian points. By avoiding the additional reconstruction errors caused by first quantizing features in the 2D image plane, our approach reduces overall reconstruction error and more efficiently utilizes the codebook in 3D space. 2) Eliminating the MLP Bottleneck: LEGaussians rely on an MLP to project each Gaussian point's features into a one-hot class category representing the index of the codebook. Our method removes the need for an MLP entirely, thereby eliminating the speed bottleneck and improving computational efficiency. 3) Efficient Sparse Coefficient Splatting: LangSplatV2 introduces Efficient Sparse Coefficient Splatting, which reduces the effective rendering dimensions to K . In contrast, LEGaussians renders with the full feature dimensions of Gaussian points, which inherently slows down the rendering process.

Furthermore, due to the reliance on an MLP decoder, LEGaussians still operate in a lower-dimensional feature space, which limits the expressiveness of its learned representations. Additionally, its two-stage 2D codebook approach fails to effectively incorporate 3D priors, resulting in a suboptimal learned field. In contrast, LangSplatV2 directly models high-dimensional features within a global 3D codebook, effectively capturing the underlying 3D structures. As evidenced by the results in Table 8, our approach outperforms LEGaussians across multiple benchmarks while being significantly faster, demonstrating superior scene reconstruction quality and more accurate language-grounded 3D representations.

To conclude, LangSplatV2 demonstrates a significant leap forward in the efficiency and scalability of modeling high-dimensional language features in 3D scenes. By leveraging the proposed 3D sparse coefficient fields and efficient sparse coefficient splatting techniques, our method reduces computational overhead while maintaining high fidelity in 3D language scene rendering.

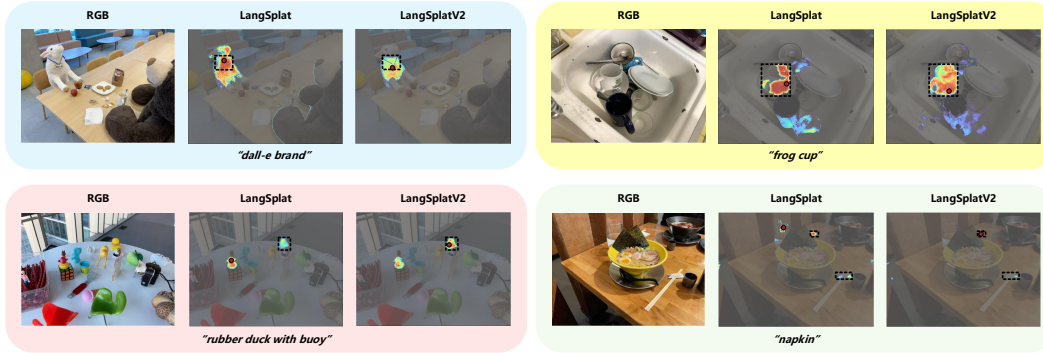


Figure 6: More qualitative comparisons of open-vocabulary 3D object localization on the LERF and Mip-NeRF360 datasets. The red points are the model predictions and the black dashed bounding boxes denote the annotations. We observe that LangSplatV2 generates better results than LangSplat.

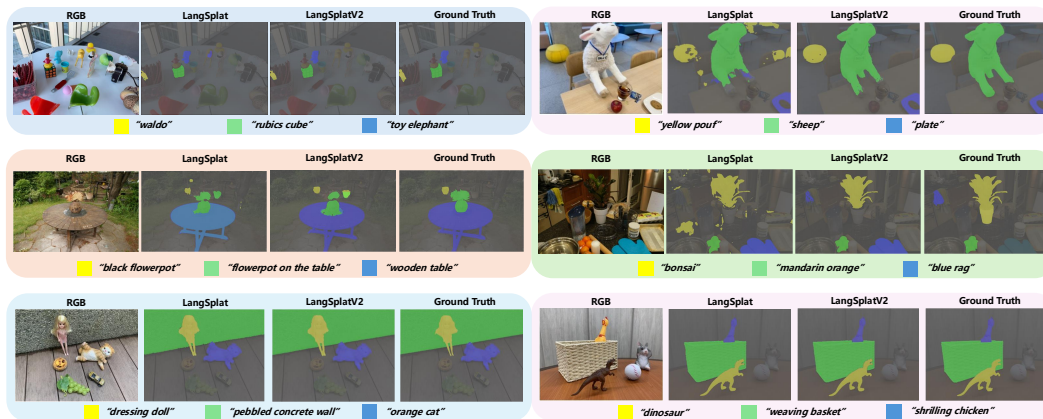


Figure 7: More qualitative comparisons of open-vocabulary 3D semantic segmentation on the LERF, Mip-NeRF360 and 3D-OVS dataset. We can see that our LangSplatV2 generates better masks than LangSplat, which shows the effectiveness of our LangSplatV2.

C Evaluation Details

After calculating the similarity between the rendered semantic features and the CLIP features of the query text, we obtained relevance maps M at different semantic levels. Then, we use different strategies for different datasets to obtain the corresponding query results.

LERF. For object localization, we first apply average pooling to smooth M and get $Avg(M)$. Then we select the point with the highest relevancy in $Avg(M)$ as the predicted point. Additionally, we choose the semantic level corresponding to the highest relevancy value among the three levels as the selected level. For semantic segmentation, after getting $Avg(M)$, we select the level in the same way as object localization. We then normalize $Avg(M)$ and use a threshold to obtain a binary mask as the final prediction mask.

3D-OVS. The post-processing pipeline for semantic segmentation is similar to that used for the LERF dataset. Differently, we obtain the smoothed relevancy map $Avg'(M)$ as $Avg'(M) = 0.5 \times (M + Avg(M))$, and the level selection is based on the average relevancy within the predicted mask. Additionally, when evaluating query accuracy, we use two coarse-grained semantic levels (*whole* and *part*) owing to the simplicity of the scenes in this dataset.

Mip-NeRF360. Similar to the post-processing for the 3D-OVS dataset, we obtain the smoothed relevancy map $Avg'(M)$ and select the semantic level based on the average relevancy score. When

evaluating query accuracy, we use all three semantic levels, as the complex scenes in this dataset demand finer-grained language feature rendering.

D More Ablation Study

Time Analysis. We conduct ablation studies on the LERF [13] dataset and report the render and decode speed (ms per query) in Table 9. We test the speed on one A100 GPU. The baseline is LangSplat, which renders three-level 3-dimensional language features separately. Rendering three semantic levels in parallel can reduce the rendering time from 6.0 ms to 2.0 ms per query. The sparse coefficient field can significantly speed up the decoding stage from 83.1 ms/q to 0.1 ms/q by changing the heavy weight decoder to a simple matrix multiplication operation. Efficient sparse coefficient splatting method effectively transformed the 192-dimensional rendering into 12-dimensional rendering, reducing the rendering time from 5.3 ms to 2.0 ms per query through CUDA optimization.

Table 9: Speed ablation study on the LERF dataset. Parallel means rendering three semantic levels in parallel, Sparse means 3D sparse coefficient field, Efficient means efficient sparse coefficient splatting with CUDA optimization.

Component			Speed (ms/q)		
Parallel	Sparse	Efficient	render	decode	total
			6.0	83.1	89.1
✓			2.0	83.1	85.1
✓	✓		5.3	0.1	5.4
✓	✓	✓	2.0	0.1	2.1

L and K . In Table 10 and Table 11, we show the ablation study results on each scene in the LERF [13] dataset. As we can see, for some complex scenarios, such as Kitchen and Figurines, larger L and K can lead to better performance. However, for all scenarios, the model’s performance is already good enough with the settings of $L = 64$ and $K = 4$, and increasing L and K will lead to increased computational resources and time cost for training and inference. Therefore, we set $L = 64$ and $K = 4$ to balance performance and efficiency.

Table 10: Ablation study on codebook size L . We report the localization and the segmentation performance on the LERF dataset.

		L	32	64	128
Ramen	Accuracy (%)		70.4	74.7	<u>74.7</u>
	IoU (%)		50.5	51.8	<u>51.4</u>
Teatime	Accuracy (%)		84.8	93.2	<u>91.5</u>
	IoU (%)		65.7	72.2	<u>69.8</u>
Kitchen	Accuracy (%)		68.2	<u>86.4</u>	86.4
	IoU (%)		54.2	<u>59.1</u>	63.2
Figurines	Accuracy (%)		67.9	<u>82.1</u>	83.9
	IoU (%)		45.3	<u>56.4</u>	57.6
Overall	Accuracy (%)		72.8	84.1	<u>84.1</u>
	IoU (%)		53.9	<u>59.9</u>	60.5

E More Visualization Results

Open-vocabulary 3D Object Localization. We visualize more examples on the LERF [13] and Mip-NeRF360 [13] datasets for open-vocabulary 3D object localization in Figure 6.

Table 11: Ablation study on different K . We report the localization and the segmentation performance on the LERF dataset.

K		2	4	6	8
Ramen	Accuracy (%)	71.8	74.7	71.8	<u>73.2</u>
	IoU (%)	50.2	51.8	<u>51.4</u>	51.1
Teatime	Accuracy (%)	91.5	93.2	<u>91.5</u>	91.5
	IoU (%)	68.8	72.2	<u>70.7</u>	70.6
Kitchen	Accuracy (%)	77.3	86.4	95.5	<u>90.9</u>
	IoU (%)	48.9	59.1	<u>59.8</u>	60.1
Figurines	Accuracy (%)	76.8	82.1	<u>78.6</u>	78.6
	IoU (%)	49.6	56.4	57.7	<u>57.6</u>
Overall	Accuracy (%)	79.4	<u>84.1</u>	84.4	83.6
	IoU (%)	54.4	59.9	<u>59.9</u>	59.9

Open-vocabulary 3D Semantic Segmentation. We visualize more examples on the LERF [13], Mip-NeRF360 [66] and 3D-OVS [47] datasets for open-vocabulary 3D semantic segmentation in Figure 7.