

GLORIE-SLAM: Globally Optimized RGB-only Implicit Encoding Point Cloud SLAM

Ganlin Zhang^{1*}, Erik Sandström^{1*}, Youmin Zhang^{2,3}, Manthan Patel¹,
Luc Van Gool^{1,4,5}, and Martin R. Oswald^{1,6}

¹ ETH Zürich, ² University of Bologna, ³ Rock Universe, ⁴ KU Leuven,
⁵ INSAIT, ⁶ University of Amsterdam

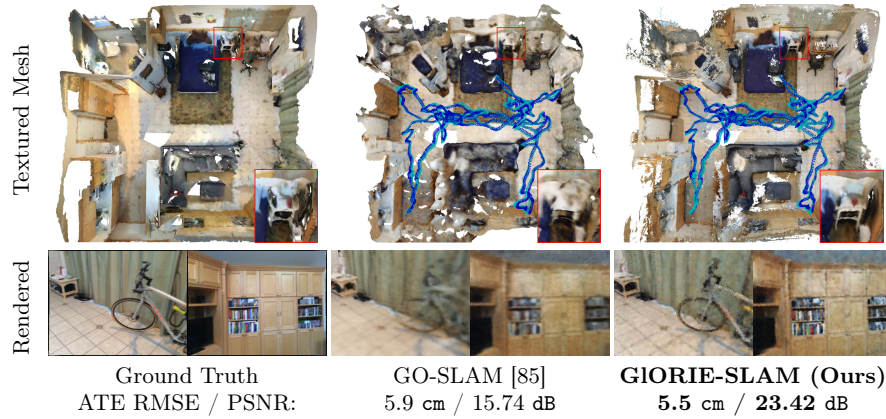


Fig. 1: GLORIE-SLAM Results on ScanNet scene 0000. GLORIE-SLAM uses a deformable point cloud as the scene representation and achieves lower trajectory error and higher rendering accuracy compared to competitive approaches, *e.g.* GO-SLAM. The geometric accuracy is qualitatively evaluated. The light blue trajectory is ground truth and the blue is the estimated. The PSNR is evaluated for all keyframes.

Abstract. Recent advancements in RGB-only dense Simultaneous Localization and Mapping (SLAM) have predominantly utilized grid-based neural implicit encodings and/or struggle to efficiently realize global map and pose consistency. To this end, we propose an efficient RGB-only dense SLAM system using a flexible neural point cloud scene representation that adapts to keyframe poses and depth updates, without needing costly backpropagation. Another critical challenge of RGB-only SLAM is the lack of geometric priors. To alleviate this issue, with the aid of a monocular depth estimator, we introduce a novel DSPO layer for bundle adjustment which optimizes the pose and depth of keyframes along with the scale of the monocular depth. Finally, our system benefits from loop closure and online global bundle adjustment and performs either better or competitive to existing dense neural RGB SLAM methods in tracking, mapping and rendering accuracy on the Replica, TUM-RGBD and ScanNet datasets. The source code will be made available.

* The two authors contributed equally to this paper.

1 Introduction

The past few years have witnessed a tremendous growth of dense SLAM, fueled predominantly by ideas such as extrapolation capability in occluded regions [62, 89], joint geometric and semantic 3D prediction [31, 87], using learnable priors for increased tracking and mapping robustness [46, 84, 86, 88], improving compression of the 3D scene representation [36, 55, 62, 79, 89] and online uncertainty estimation [54, 56]. A common factor for the majority of these works is that they reconstruct the dense map by optimizing a neural implicit encoding of the scene in the form of either weights of an MLP [1, 42, 50, 62] or as features anchored in dense grids [3, 30, 47, 56, 63, 73, 74, 89, 90], hierarchical octrees [79], via voxel hashing [8, 43, 54, 84, 85], point clouds [18, 33, 55] or axis aligned feature planes [36, 51]. Concurrent to our work, we have also seen the introduction of 3D Gaussian Splatting (3DGS) to the dense SLAM field [22, 26, 41, 77, 81].

When surveying the recent trends, we make certain discoveries: **1.** Explicit representations such as point clouds achieve state-of-the-art rendering and reconstruction accuracy [55]. **2.** The majority of works perform frame-to-model tracking and struggle to efficiently introduce bundle adjustment (BA) and/or loop closure [19, 29, 32, 36, 55, 56, 71, 76, 88]. Instead, the methods with the lowest camera pose drift on real-world data are based on external trackers such as DROID-SLAM [67], extended with online loop closure and global BA [84, 85], or ORB-SLAM2 [8]. **3.** The RGB-only setting is under-explored compared to the RGB-D setting as it needs to deal with depth and scale ambiguities, making the problem more difficult. Recently, HI-SLAM [84] and GO-SLAM [85] introduce dense RGB-only SLAM frameworks. They use hierarchical feature grids as the dense map representation, which does not efficiently and accurately lend itself for map deformations. Such map deformations are required when performing loop closure or BA to avoid recomputing the map from scratch and hence requiring to save all input data as commonly done in the literature [23, 54].

To this end, we propose to use an efficient and accurate neural point cloud scene representation to construct an RGB-only SLAM system that can also handle complex and large scale indoor scenes by integrating online loop closure along with online global bundle adjustment (BA) and a monocular depth prior which aids not only in reconstruction completeness, but also in accuracy.

In summary, our **contributions** include:

- An RGB-only dense SLAM framework which performs camera pose estimation via dense optical flow tracking and includes loop closure and online global BA. Mapping is done with a neural point cloud map representation that allows for online rigid map deformations (instead of costly map re-creation) for online loop closures as well as pose and geometry refinement. This results in more accurate reconstructions compared to hierarchical grid-based approaches while not requiring backpropagation to retrain the grid-anchored features. See Fig. 1.
- A novel Disparity, Scale and Pose Optimization (DSPO) layer that combines pose and geometry estimation. It refines inaccurate parts of the estimated keyframe disparity by tightly coupling a monocular depth prior into the bundle adjustment.

2 Related Work

Dense Visual SLAM. The pioneering work of Curless and Levoy [9] using truncated signed distance functions (TSDF) laid the foundation for dense online 3D reconstruction strategies. KinectFusion [47] demonstrated real-time dense mapping and tracking using depth maps and scalability was further improved through techniques like voxel hashing [11, 25, 43, 48, 49], and octrees [5, 34, 40, 58, 79]. Besides grid based methods, point based SLAM has also been successful [4, 6, 25, 27, 55, 57, 75, 82], typically in the form of surface elements (surfels). To obtain globally consistent pose estimates and dense maps, numerous techniques have been developed. The most popular includes splitting the global map into parts, usually called submaps [2, 4, 7, 11, 15, 16, 24, 25, 37–39, 43, 53, 60, 65, 65], which aggregate the information from a set of frames. Once a loop is detected, the submaps are rigidly registered via pose graph optimization [4, 7, 13, 14, 16–18, 24, 28, 37, 39, 43, 43, 53, 57, 60, 65, 70, 78]. To refine the solution, global BA is sometimes applied [4, 8, 11, 18, 43, 57, 65, 67, 78, 80].

Most similar to our work, but in the RGBD setting, is the concurrent Loopy-SLAM [33], which also utilizes the map representation from Point-SLAM [55]. They split the global map into submaps and apply a robust pose graph formulation based on direct point cloud registration. Other concurrent RGBD works include methods that use 3D Gaussians as the scene representation [26, 77, 81]. None of these works consider global map consistency via *e.g.* loop closure. For a recent in depth survey on NeRF inspired dense SLAM, we refer to [69].

RGB-only Dense Visual SLAM. The majority of research focuses on combining RGB and depth cameras, while methods using only RGB cameras are less explored due to their lack of geometric data, causing scale and depth ambiguities. However, RGB-only dense SLAM is appealing for its cost-effectiveness and versatility in various environments. Within the scope of NeRF inspired dense SLAM, NeRF-SLAM [54] and Orbeez-SLAM [8] use DROID-SLAM [67] and ORB-SLAM2 [45] as the tracking modules respectively. Both use Instant-NGP [44] as the volumetric neural radiance field map. NeRF-SLAM does not tackle the global consistency of the map and Orbeez-SLAM [8] requires undesirable training iterations of the hash features to perform map corrections. DIM-SLAM [29] and NICER-SLAM [88] both perform tracking and mapping on the same neural implicit map representation, consisting of hierarchical feature grids and do not address global map consistency via *e.g.* loop closure. Recently, GO-SLAM [85] and Hi-SLAM [84] adapts DROID-SLAM [67] to the full SLAM setting by introducing online loop closure via factor graph optimization. Both systems achieve high tracking accuracy, but like Orbeez-SLAM and NeRF-SLAM, they require backpropagation steps to update the feature encodings in the Instant-NGP map. We are inspired by both works, and make some important modifications to increase tracking and mapping accuracy. We integrate our own point based map representation to account for online updates of poses and geometry as a result of bundle adjustment (BA) or loop closure.

Concurrent to our work, FMapping [21], Hi-Map [20], TT-HO-SLAM [46] and NeRF-VO [46] all utilize hierarchical grid based data structures. MoD-SLAM

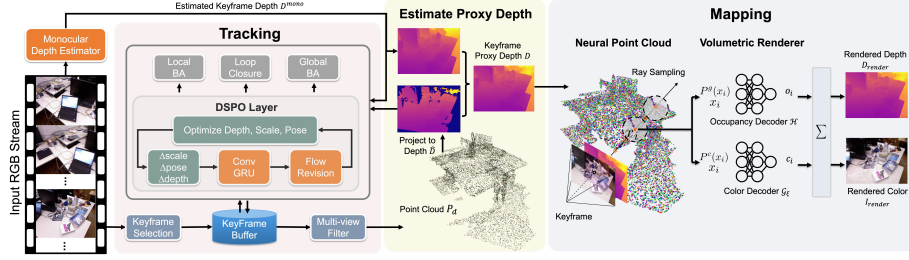


Fig. 2: GLORIE-SLAM Architecture. Given an input RGB stream, we first track and then map every keyframe. The pose is initially estimated with local bundle adjustment (BA) via frame-to-frame tracking of recurrent optical flow estimation. This is done with our novel DSPO (Disparity, Scale and Pose Optimization) layer, which combines pose and depth estimation with scale and depth refinement by leveraging a monocular depth prior. The DSPO layer also refines the poses globally via online loop closure and global BA. To map the estimated pose, a proxy depth map is estimated by combining the noisy keyframe depths from the tracking module with the monocular depth prior to account for missing observations. Mapping is done, along with the input RGB keyframe via a deformable neural point cloud, leveraging depth guided volumetric rendering. A re-rendering loss to the input RGB and proxy depth optimizes the neural features and the color decoder weights. Importantly, the neural point cloud deforms to account for global updates of the poses and proxy depth before each mapping phase.

uses an MLP to parameterize the map via a unique reparameterization. Perhaps most closely related to our work are MonoGS [41] and Photo-SLAM [22]. They both tackle monocular dense SLAM by modeling the map with 3D Gaussians.

Depth Priors for RGB-only SLAM. NICER-SLAM [88] estimates the scale and shift of a relative mono-depth estimator and supervises all pixels equally. MoD-SLAM [86] combines relative and metric mono-depth estimation, and requires additional finetuning of the metric depth. HI-SLAM [84] proposes a similar technique to ours, but regularizes all available keyframe depth pixels with the mono-depth prior. In our DSPO layer, we instead split the optimization and use the monocular prior to regularize the high error keyframe depth pixels while the low error keyframe depth is kept fixed to stabilize scale estimation.

3 Method

As an RGB-only SLAM framework, our aim is to simultaneously track the camera poses and reconstruct globally consistent scene geometry. To achieve this, we utilize a deformable neural point cloud as the scene representation (Sec. 3.1). During mapping, we render depth and color from the neural point cloud (Sec. 3.2), and optimize a re-rendering loss with respect to the sensor input (Sec. 3.3). For tracking, we use an optical-flow-based method containing loop closure and online global bundle adjustment. The tracker takes an RGB stream along with relative depth from a monocular depth estimator as input and outputs is a set of estimated camera poses and noisy semi-dense depth maps and feeds them to the mapper (Sec. 3.4). For an overview of our method, see Fig. 2.

3.1 Neural Point Cloud Scene Representation

Similar to Point-SLAM [55], we use a neural point cloud to represent the scene, but the RGB-only setting requires several modifications. The neural point cloud with N neural points is defined as

$$P = \{(p_i, f_i^g, f_i^c, k_i, u_i, v_i, D_i) \mid i = 1, \dots, N\} , \quad (1)$$

with point locations $p_i \in \mathbb{R}^3$ and $f_i^g, f_i^c \in \mathbb{R}^{32}$ being the geometric and color features encoding the local shape and appearance information of p_i and its surrounding. The other attributes are used to deform the map at *e.g.* loop closure. k_i is the frame index that anchored point i from pixel (u_i, v_i) at depth D_i .

Point Adding Strategy. For every keyframe, given the estimated camera pose and proxy depth map (see Sec. 3.2), we adopt the point adding strategy from Point-SLAM *i.e.* we sample X pixels uniformly and Y pixels among the top $5Y$ pixels according to the color gradient magnitude. This ensures an even spread of pixels while focusing on reconstructing complex areas properly. The $X + Y$ pixels are projected into 3D space, and if no existing point is found within radius r , three neural points are initialized along the ray, with depth $(1 - \rho)D$, D and $(1 + \rho)D$ respectively, with D being the proxy depth of the pixel and $\rho \in (0, 1)$ is a hyperparameter. The search radius r is computed dynamically as a linear transformation of the pixel color gradient $\nabla I(u, v)$, and clamped by the predefined maximum and minimum radius r_u and r_l respectively. Different from Point-SLAM, the scale of the scene cannot be determined a priori in the RGB-only setting. We therefore scale the radius with the estimated depth to adapt to the reconstructed scene size as

$$r(u, v) = D(u, v) \max \left(\min(\beta_1 \nabla I(u, v) + \beta_2, r_u), r_l \right) , \quad (2)$$

where $\beta_1, \beta_2 \in \mathbb{R}$ are hyperparameters.

Map Update. In a SLAM system, the estimated set of camera poses are globally updated as a result of bundle adjustment and loop closure. For this reason, the map needs to be deformed to ensure it is consistent with the updated poses. Thanks to our point-based representation, map deformations are straightforward. For every tuple of three points, which are added along a ray in the neural point cloud, we update the locations p_i by rigidly re-anchoring them. Specifically, for the central point, the new point location p_i' is computed as,

$$p_i' = \omega_{k_i}' D_i' K^{-1} [u_i, v_i, 1]^T, \quad (3)$$

where $\omega_{k_i}' \in \mathbb{SE}(3)$ is the updated camera-to-world pose, D_i' the updated depth and K the temporally constant camera intrinsics. Note that only the locations of the points are updated and not the neural features. For points that need to be deformed, but lack updated depth D_i' , we rescale the depth map that first initialized and anchored the neural point. We do this with a least squares fitting to the new depth map such that $D_i' = s D_i$, where s is the rescaling factor and D_i is the depth that was available when the point was initialized.

3.2 Rendering

Similar to Point-SLAM [55], we use depth-guided volume rendering that sparsely samples the 3D space for faster RGB and depth rendering.

Proxy Depth Map. Different from Point-SLAM, which benefits from a dense depth sensor, the estimated monocular depth is less reliable. To still benefit from sparse volume rendering, we construct a proxy depth map for each keyframe. We utilize the estimated dense but noisy depth map $\tilde{D}$ from the tracker and further filter out inconsistent depth with multi-view constraints as follows. For a given depth $\tilde{D}_c$ from keyframe c , we project every pixel (u, v) to 3D as p_c , which is then projected into keyframe j to obtain the corresponding pixel $(\hat{u}, \hat{v})$ as

$$p_c = \omega_c \tilde{D}_c(u, v) K^{-1} [u, v, 1]^T, \quad [\hat{u}, \hat{v}, 1]^T \propto K \omega_j^{-1} [p_c, 1]^T. \quad (4)$$

We unproject the depth at $(\hat{u}, \hat{v})$ to obtain the corresponding 3D location p_j ,

$$p_j = \omega_j \tilde{D}_j(\hat{u}, \hat{v}) K^{-1} [\hat{u}, \hat{v}, 1]^T. \quad (5)$$

If the L2 distance between p_c and p_j is small enough, the estimated depth $\tilde{D}_c(u, v)$ is consistent between these two views. The global two-view consistency n_c for frame c can be computed by looping over all keyframes as

$$n_c(u, v) = \sum_{\substack{k \in \text{KFs}, \\ k \neq c}} \mathbf{1} \left(\|p_c - p_k\| < \eta \cdot \text{average}(\tilde{D}_c) \right), \quad (6)$$

where $\mathbf{1}(\cdot)$ is the indicator function. Here, $\eta \in \mathbb{R}_{\geq 0}$ is a hyperparameter and n_c is the total two-view consistency for pixel (u, v) in keyframe c . If n_c is larger than a threshold, $\tilde{D}_c(u, v)$ is valid. We form a point cloud P_d by unprojecting all depth values from the keyframes to 3D.

$$P_d = \left\{ \omega_k \tilde{D}_k(u, v) K^{-1} [u, v, 1]^T \mid k \in \text{KFs}, \tilde{D}_k(u, v) \text{ is valid} \right\}. \quad (7)$$

Next, we utilize the filtered multi-view information in P_d and form almost dense depth maps $\hat{D}_c$ by projecting P_d into the keyframes. The more frames that are observed, the denser they become. Nevertheless, $\hat{D}_c$ can still have missing pixels. To fill these, we use a monocular depth prior D_c^{mono} predicted by an off-the-shelf mono-depth estimator [12]. The predicted depth is, however, only estimated up to unknown scale θ_c and shift γ_c . These parameters are estimated by fitting the monocular depth to the depth maps $\hat{D}_c$ with least squares as

$$\hat{\theta}_c, \hat{\gamma}_c = \arg \min_{\theta, \gamma} \sum_{(u, v)} \left((\theta D_c^{\text{mono}}(u, v) + \gamma) - \hat{D}_c(u, v) \right)^2. \quad (8)$$

Then, the final proxy depth map D_c is,

$$D_c(u, v) = \begin{cases} \hat{D}_c(u, v) & \text{if } \hat{D}_c(u, v) \text{ has value} \\ \hat{\theta}_c D_c^{\text{mono}}(u, v) + \hat{\gamma}_c & \text{otherwise} \end{cases}. \quad (9)$$

Rendering RGB and Depth. Given the estimated camera pose and corresponding proxy depth map, we sample a set of 3D points along a ray,

$$p_i = o_c + z_i v_i , \quad (10)$$

where $o_c \in \mathbb{R}^3$ is the camera origin, $z_i \in \mathbb{R}$ the point depth and $v_i \in \mathbb{R}^3$ the ray direction. We sample 10 points spread evenly between $(1-\rho)D_c$ and $(1+\rho)D_c$ (D_c is the proxy depth) to ensure the real surface is contained within the extraction band, as the proxy depth can be noisy. For each sampled point p_i , two MLPs $\mathcal{H}$ and $\mathcal{G}_\xi$ decode the occupancy $s_i \in [0, 1]$ and color $t_i \in \mathbb{R}^3$ respectively as

$$s_i = \mathcal{H}(p_i, P^g(p_i)) , \quad t_i = \mathcal{G}_\xi(p_i, P^c(p_i), v_i) . \quad (11)$$

For the occupancy decoder $\mathcal{H}$, we use the same MLP as Point-SLAM [55] and use their pretrained and fixed parameters. For the color decoder $\mathcal{G}_\xi$, we modify the MLP used by [55] by additionally feeding the view direction v_i as input. The model parameters ξ are optimized on the fly. To obtain the geometric feature P^g and color feature P^c at the sampled points p_i , we search neighboring neural points within the query radius $2r$ in the neural point cloud and use a distance-based weighted average of the features f^g and f^c of these neural points. We refer to [55] for the details of the point cloud feature interpolation. The per-point occupancy and color are used to render the per-pixel depth and color via volume rendering. Specifically, we render the depth D_{render} and color I_{render} as the weighted average of the depth and color values along each ray,

$$D_{\text{render}} = \sum_{i=1}^N \alpha_i z_i , \quad I_{\text{render}} = \sum_{i=1}^N \alpha_i t_i , \quad (12)$$

where $\alpha_i = s_i \prod_{j=i}^{i-1} (1 - s_j)$ can be interpreted as the probability that the ray terminates at point p_i , and N is the number of sampled points along each ray.

3.3 Loss Function

For each mapping phase, we first randomly select κ keyframes which overlap significantly with the viewing frustum of the current frame. Among the selected keyframes including the current frame, we sample M pixels uniformly across the image plane similar to [55]. Rendering is done as described in Sec. 3.2 and we calculate the L_1 loss using the sensor RGB image and the proxy depth map,

$$\mathcal{L}_{\text{geo}} = \sum_{m=1}^M |D_m - D_m^{\text{render}}|_1 , \quad \mathcal{L}_{\text{color}} = \sum_{m=1}^M |I_m - I_m^{\text{render}}|_1 . \quad (13)$$

We add an additional frame-to-frame pixel warping term to ensure geometric consistency of the map from different views. For each sampled pixel (u, v) in keyframe c , we first project it to 3D as p_m , using the rendered depth D_{render} as in

Eq. (5). Then p_m is projected into other selected keyframes and the photometric loss between the projected pixel $(\hat{u}, \hat{v})$ and the original pixel is applied:

$$\mathcal{L}_{\text{pix}} = \sum_{m=1}^M \sum_{\substack{k \in \text{KFs}, \\ k \neq c}} |I_c(u, v) - I_k(\hat{u}, \hat{v})|_1 . \quad (14)$$

The final loss is

$$\mathcal{L} = \lambda_{\text{geo}} \mathcal{L}_{\text{geo}} + \lambda_{\text{pix}} \mathcal{L}_{\text{pix}} + \lambda_{\text{color}} \mathcal{L}_{\text{color}} , \quad (15)$$

with hyperparameters $\lambda_{\text{geo}}, \lambda_{\text{pix}}, \lambda_{\text{color}} \in \mathbb{R}_{\geq 0}$. Importantly, we split the optimization into two stages. For the first 30 % of iterations, $\lambda_{\text{color}} = 0$ to initialize the geometry well. Then, $\lambda_{\text{color}} > 0$ and f^g, f^c and ξ are jointly optimized.

3.4 Tracking

To robustly track the camera under challenging scenarios, *e.g.* repeated patterns, low texture or rapid camera movement, we adopt the idea of frame-to-frame tracking [85] which is based on recurrent optical flow [66]. A factor graph $\mathcal{G}(\mathcal{V}, \mathcal{E})$ is maintained during the tracking process, where every node in $\mathcal{V}$ stores the pose and estimated disparity of a keyframe and every edge in $\mathcal{E}$ stores the optical flow between two keyframes. For every incoming frame, we calculate the flow with respect to the last added keyframe in the graph. If the mean flow magnitude is larger than a threshold $\tau \in \mathbb{R}$, it is added to the factor graph and optimized.

We use the Dense Bundle Adjustment (DBA) layer from [67] to optimize the pose and disparity of the keyframes, as shown in Eq. (16). This is performed in a sliding window fashion over a local window overlapping the current keyframe and is equivalent to tracking via local BA, as both the camera pose and the geometry are estimated jointly.

$$\arg \min_{\omega, d} \sum_{(i,j) \in \mathcal{E}} \left\| \tilde{p}_{ij} - K \omega_j^{-1} (\omega_i (1/d_i) K^{-1} [p_i, 1]^T) \right\|_{\Sigma_{ij}}^2 , \quad (16)$$

In Eq. (16), $\tilde{p}_{ij} \in \mathbb{R}^{(W \times H \times 2) \times 1}$ denotes the flattened predicted pixel coordinates when the pixel grid $p_i \in \mathbb{R}^{(W \times H \times 2) \times 1}$ from keyframe i is projected into keyframe j using optical flow. $d_i \in \mathbb{R}^{(W \times H) \times 1}$ is the disparity map of keyframe i , K the camera intrinsics and ω_j and ω_i the camera to world extrinsics for keyframes j and i respectively. Finally, $\|\cdot\|_{\Sigma_{ij}}$ denotes the Mahalanobis distance, with Σ_{ij} being a diagonal matrix composed of the prediction confidences of the optical flow for each flow term in $\tilde{p}_{ij}$. Notice that we make two simplifications of the equation: 1. We do not convert the 3D points to homogeneous coordinates and 2. We do not divide by the third coordinate to convert to the pixel space.

DSPO Layer. We find that disparity maps d can still be noisy when only using the DBA layer, which may affect the scale and shift estimation negatively. Therefore, we add another optimization objective besides the DBA layer and

optimize them alternatively. We call this the Disparity, Scale and Pose Optimization (DSPO) layer which incorporates the predicted monocular depth prior D^{mono} into the tracking loop. We alternately optimize the pose and disparity in Eq. (16) and the scale θ , shift γ and the high error disparities d^h as

$$\arg \min_{d^h, \theta, \gamma} \sum_{(i,j) \in \mathcal{E}} \|\tilde{p}_{ij} - K\omega_j^{-1}(\omega_i(1/d_i^h)K^{-1}[p_i, 1]^T)\|_{\Sigma_{ij}}^2 \quad (17)$$

$$+ \alpha_1 \sum_{i \in \mathcal{V}} \|d_i^h - (\theta_i(1/D_i^{\text{mono}}) + \gamma_i)\|^2 + \alpha_2 \sum_{i \in \mathcal{V}} \|d_i^l - (\theta_i(1/D_i^{\text{mono}}) + \gamma_i)\|^2.$$

In Eq. (17), d_i^l and d_i^h are the valid (low error) and invalid (high error) parts of the disparity map for frame i as defined by the two-view consistency threshold in Eq. (6). During optimization, the high error disparity d^h is regularized by the monocular prior, while the scale and shift are refined by the low error disparity d^l . We use hyperparameters $\alpha_1 < \alpha_2$, giving more weight to pixels with accurate disparities ensuring that the scale and shift converge to the right local minima while at the same time optimizing the inaccurate part of the disparity map. The scale θ and shift γ are initialized as in Eq. (8), with the help of point cloud P_d . As also found by [84], because of scale ambiguity, we do not optimize d , θ , γ and ω together, but alternately via Eq. (16) and Eq. (17) using the Gauss-Newton algorithm. Contrary to [84], we only optimize the high-error disparity pixels, while letting the low-error pixels refine the scale θ and shift γ . We deploy the DSPO layer for local BA, loop closure and global BA (details below).

Loop Closure. Apart from frame tracking within a local window, we use loop closure and online global BA to avoid scale and pose drift. For loop detection, we compute the mean optical flow magnitude between the current active keyframes (within the local window) and all past keyframes. For every pair, two conditions are checked: First, the flow should fall below a threshold τ_{loop} to ensure there is enough co-visibility between the two views. Second, the time interval between these two frames need to be larger than a predefined threshold τ_t to avoid redundant edges in the graph. If both conditions are met, a unidirectional edge is added to the graph. During loop closure optimization, only the active keyframes and connected loop nodes are optimized to reduce the computational cost.

Global BA. For online global BA, we build an additional graph containing all keyframes so far. Edges are added based on the temporal and spatial distance between the keyframe pairs as in [85]. We found that the open sourced code of [85] did not give stable results. To alleviate this, we run online global BA every 20 keyframes. To further improve numerical stability during optimization, we normalize the scale of the disparity maps and poses before each round of global BA. This is done by computing the mean disparity $\bar{d}$ over all keyframes, and updating the disparity as $d_{\text{norm}} = d/\bar{d}$ and the pose translation t as $t_{\text{norm}} = \bar{d}t$.

4 Experiments

We first describe our experimental setup and then evaluate our method against state-of-the-art dense neural RGB and RGBD SLAM methods on Replica [59]

Metrics	GO-SLAM [85]	NICER-SLAM [88]	MoD-SLAM* [29]	Photo-SLAM* [22]	Ours
PSNR $\uparrow$	22.13	25.41	27.31	33.30	31.04
SSIM $\uparrow$	0.73	0.83	0.85	0.93	0.97
LPIPS $\downarrow$	-	0.19	-	-	0.12
ATE RMSE $\downarrow$	0.39	1.88	0.35	1.09	0.35

Table 1: Rendering and Tracking Results on Replica [59] for RGB-Methods. Our method outperforms all methods on rendering except the concurrent Photo-SLAM*, where we see similar performance. Our method achieves the best tracking accuracy among all methods. Results are averaged over 8 scenes. Results are from [69] except ours. The best results are highlighted as **first**, **second**, and **third**.

as well as the real world TUM-RGBD [61] and the ScanNet [10] datasets. For more experiments and details, we refer to the supplementary material.

Implementation Details. We use $\rho = 0.05$ for the depth interval for both anchoring neural points and rendering. For the search radius parameters, we use $\beta_1 = -0.4$, $\beta_2 = 0$, and $r_u = 0.027$, $r_l = 0.007$ are the upper-bond and lower-bond radii. For the proxy depth, we use $\eta = 0.01$ to filter points which are inconsistent between view pairs, and use the condition $n_c \geq 2$ to ensure multi-view consistency. For the mapping loss function, we use $\lambda_{\text{geo}} = 1.0$, $\lambda_{\text{pix}} = 1000.0$ and $\lambda_{\text{color}} = 0.1$. We sample $M = 2K$ pixels for each mapping phase in ScanNet and TUM-RGBD, and $M = 5K$ in Replica. For the number of selected keyframes during the mapping, we use $\kappa = 12$ in Replica, $\kappa = 10$ in TUM-RGBD and $\kappa = 5$ in ScanNet. For tracking, we use $\alpha_1 = 0.01$ and $\alpha_2 = 0.1$ as weights for the DSPO layer. We use the flow threshold $\tau = 4.0$ on ScanNet, $\tau = 3.0$ on TUM-RGBD and $\tau = 2.25$ on Replica. The threshold for loop detection is $\tau_{\text{loop}} = 25.0$. The time interval threshold is $\tau_t = 20$. We conducted the experiments on a cluster with an 1.50GHz AMD EPYC 7742 64-Core CPU and an NVIDIA GeForce RTX 3090 with 24 GiB GPU memory.

Evaluation Metrics. We evaluate three SLAM tasks: rendering and reconstruction quality as well as tracking accuracy. For rendering we report PSNR, SSIM [72] and LPIPS [83] on the rendered keyframe images against the sensor images. For reconstruction, we first extract the meshes with marching cubes [35] as in [55] and evaluate the meshes using accuracy [cm], completion [cm] and completion ratio [%] against the provided ground truth meshes. Further, the depth L1 [cm] evaluates the depth on the mesh at random poses against ground truth as in [89]. We use ATE RMSE [cm] [61] to evaluate the estimated trajectory.

Datasets. The Replica dataset [59] consists of high-quality synthetic 3D reconstructions of diverse indoor scenes. We leverage the publicly available dataset [62], which contains trajectories from an RGBD sensor. Additionally, we showcase our framework on real-world data using the TUM-RGBD dataset [61] and the ScanNet dataset [10]. The TUM-RGBD poses were captured utilizing an external motion capture system, while ScanNet uses poses from BundleFusion [11].

Baseline Methods. We compare our method to numerous published and concurrent works on dense RGB and RGBD SLAM. Concurrent works are denoted with an asterisk*. The main baselines are GO-SLAM [85] and HI-SLAM [84].

Method	Metric	0000	0059	0106	0169	0181	0207	Avg.
<i>RGB-D Input</i>								
NICE-SLAM [89]	PSNR↑	18.71	16.55	17.29	18.75	15.56	18.38	17.54
	SSIM ↑	0.641	0.605	0.646	0.629	0.562	0.646	0.621
	LPIPS↓	0.561	0.534	0.510	0.534	0.602	0.552	0.548
VOX-Fusion [79]	PSNR↑	19.06	16.38	18.46	18.69	16.75	19.66	18.17
	SSIM ↑	0.662	0.615	0.753	0.650	0.666	0.696	0.673
	LPIPS↓	0.515	0.528	0.439	0.513	0.532	0.500	0.504
ESLAM [36]	PSNR↑	15.70	14.48	15.44	14.56	14.22	17.32	15.29
	SSIM ↑	0.687	0.632	0.628	0.656	0.696	0.653	0.658
	LPIPS↓	0.449	0.450	0.529	0.486	0.482	0.534	0.488
Point-SLAM [55]	PSNR↑	21.30	19.48	16.80	18.53	22.27	20.56	19.82
	SSIM ↑	0.806	0.765	0.676	0.686	0.823	0.750	0.751
	LPIPS↓	0.485	0.499	0.544	0.542	0.471	0.544	0.514
<i>RGB Input</i>								
GO-SLAM [85]	PSNR↑	15.74	13.15	14.58	14.49	15.72	15.37	14.84
	SSIM ↑	0.423	0.315	0.459	0.416	0.531	0.385	0.421
	LPIPS↓	0.614	0.598	0.594	0.574	0.615	0.602	0.599
Ours	PSNR↑	23.42	20.66	20.41	25.23	21.28	23.68	22.45
	SSIM ↑	0.867	0.828	0.840	0.912	0.759	0.850	0.842
	LPIPS↓	0.262	0.306	0.313	0.214	0.439	0.287	0.304

Table 2: Rendering Performance on ScanNet [10]. Our method performs even better than all RGB-D methods on commonly reported rendering metrics. For all RGB-D methods, we take the numbers from [81].

Metrics	NeRF-SLAM [69]	DIM-SLAM [29]	GO-SLAM [85]	NICER-SLAM [88]	HI-SLAM [84]	MoD-SLAM* [86]	Ours
Depth L1 ↓	4.49	-	4.39	-	3.63	3.23	3.24
Accuracy ↓	-	4.03	3.81	3.65	3.62	2.48	2.96
Completion ↓	-	4.20	4.79	4.16	4.59	-	3.95
Comp. Rat. ↑	-	79.60	78.00	79.37	80.60	-	83.72

Table 3: Reconstruction Results on Replica [59] for RGB-Methods. Our method achieves the best reconstruction performance on all metrics compared to published works. We perform similarly to the concurrent MoD-SLAM*, which finetunes a mono-depth estimator, while ours does not. Results are averaged over 8 scenes.

Rendering. In Tab. 1, we evaluate the rendering performance on Replica [59]. We outperform all published works, while beating the concurrent MoD-SLAM* [86] and being competitive with Photo-SLAM* [22], based on Gaussian splatting. A qualitative evaluation is shown in Fig. 3. Compared to GO-SLAM [85], our method does not suffer from blurry artifacts and can even render high-frequency details well, *e.g.* the blinds in the first and second column. We further show competitive performance on the real-world dataset ScanNet [10] in Tab. 2. Our method performs the best, including RGB-D methods. We attribute this to our deformable point cloud scene representation, which allows to be updated without further optimization when the estimated camera trajectory is globally refined. None of the RGB-D methods employ loop closure and suffer from pose drift, resulting in worse rendering performance. GO-SLAM, which is globally optimized requires feature refinement every time a global update is done, meaning that the optimization never settles, leading to the worst performance on all metrics.

Reconstruction. We report reconstruction metrics on Replica in Tab. 3. Our method achieves the best performance on completion and completion ratio,

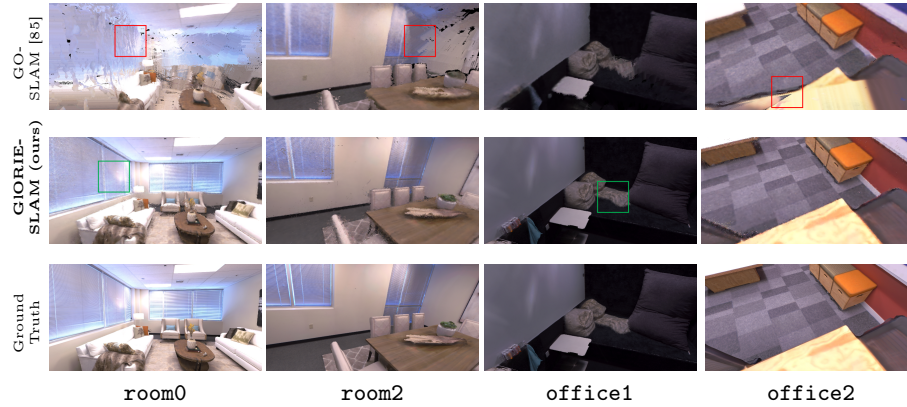


Fig. 3: Rendering Results on Replica [59]. The red boxes show blurry artifacts from GO-SLAM because of insufficient optimization when camera poses and depth are updated. Ours does not suffer from that by deforming the points accordingly. The green boxes show that ours can render high-frequency details well.

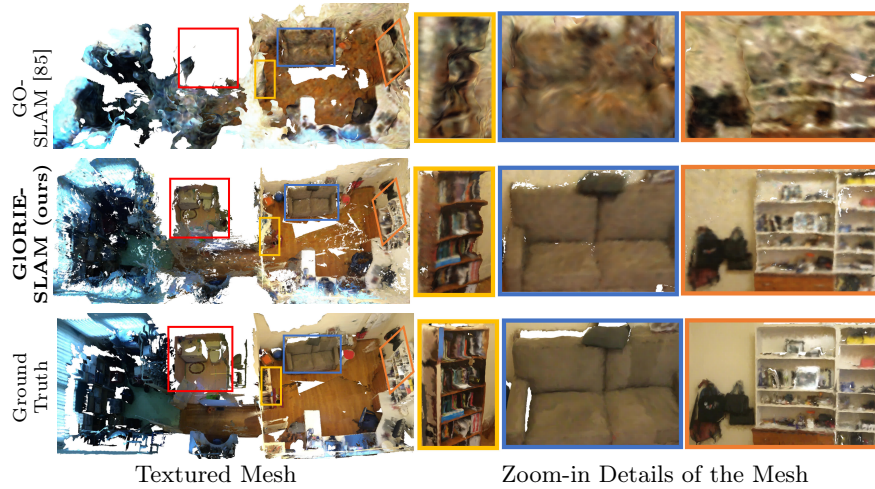


Fig. 4: Reconstruction Results on ScanNet scene 0054. The red box shows the bathroom which GO-SLAM fails to reconstruct while ours succeeds. The yellow, blue and orange boxes show that our method can produce more detailed reconstructions.

for depth L1 and accuracy ours also achieves competitive results with MoD-SLAM* [86], which finetunes mono-depth estimator and does extra training while ours does not. Fig. 4 shows qualitative results on a ScanNet multi-room scene.

Tracking. In Tab. 1, Tab. 4 and Tab. 5, we report the tracking accuracy of the estimated trajectory on Replica [59], ScanNet [10] and TUM-RGBD [61]. On all three datasets, our method shows competitive results in every single scene and gives the best average value among the RGB and RGB-D methods.

Ablation Study. In Tab. 6, we conduct a set of ablation studies related to our method, by enabling and disabling certain parts. We find that deforming

Method	0000	0059	0106	0169	0181	0207	Avg.-6	0054	0233	Avg.-8
<i>RGB-D Input</i>										
NICE-SLAM [89]	12.0	14.0	7.9	10.9	13.4	6.2	10.7	20.9	9.0	11.8
Co-SLAM [71]	7.1	11.1	9.4	5.9	11.8	7.1	8.7	-	-	-
ESLAM [36]	7.3	8.5	7.5	6.5	9.0	5.7	7.4	36.3	4.3	10.6
<i>RGB Input</i>										
GO-SLAM [85]	5.9	8.3	8.1	8.4	8.3	6.9	7.7	13.3	5.3	8.1
HI-SLAM [84]	6.4	7.2	6.5	8.5	7.6	8.4	7.4	-	-	-
Ours	5.5	9.1	7.0	8.2	8.3	7.5	7.6	9.4	5.1	7.5

Table 4: Tracking Accuracy ATE RMSE [cm] ↓ on ScanNet [10]. Our method performs on average competitively with HI-SLAM and better than all other methods. Results for the RGB-D methods are from [33].

Method	f1/desk	f2/xyz	f3/off	Avg.-3	f1/desk2	f1/room	Avg.-5
<i>RGB-D Input</i>							
SplaTAM* [26]	3.4	1.2	5.2	3.3	6.5	11.1	5.5
GS-SLAM* [77]	1.5	1.6	1.7	1.6	-	-	-
GO-SLAM [85]	1.5	0.6	1.3	1.1	-	4.7	-
<i>RGB Input</i>							
MonoGS* [41]	4.2	4.8	4.4	4.4	-	-	-
Photo-SLAM* [22]	1.5	1.0	1.3	1.3	-	-	-
DIM-SLAM [29]	2.0	0.6	2.3	1.6	-	-	-
GO-SLAM [85]	1.6	0.6	1.5	1.2	2.8	5.2	2.3
MoD-SLAM* [86]	1.5	0.7	1.1	1.1	-	-	-
Ours	1.6	0.2	1.4	1.1	2.8	4.2	2.1

Table 5: Tracking Accuracy ATE RMSE [cm] ↓ on TUM-RGBD [61]. Our method performs on average the best among all the methods. Note that all methods with a * are concurrent works.

the neural point cloud is important for rendering and geometry prediction. The monocular depth prior is important for the geometric completeness and the DSPO layer helps improve both the rendering and reconstruction accuracy. In Fig. 5, we compare the valid estimated depth $\tilde{D}$, showing that our proposed DSPO layer contributes significantly to more consistent depth estimation.

Memory and Running Time. In Tab. 7, we evaluate the peak GPU memory usage and runtime of our method. We achieve a comparable memory usage with GO-SLAM [85] which is the most representative method using frame-to-frame tracking. SplaTAM [26], which is a RGB-D 3D GS method has a similar memory footprint. We use more memory compared to the RGB-D method

Deformable Pointcloud	Mono Prior in Proxy Depth	DSPO Layer	PSNR [dB] ↑	Depth L1 [cm] ↓	Acc. [cm] ↓	Comp. [cm] ↓	Comp. Ratio [%] ↑
✗	✗	✗	25.52	9.21	5.91	7.70	73.02
✗	✓	✗	25.58	9.07	5.76	6.61	76.78
✓	✗	✗	30.89	3.55	3.61	5.19	78.84
✓	✓	✗	30.66	3.52	3.72	3.92	83.40
✓	✓	✓	31.04	3.24	2.96	3.95	83.72

Table 6: Ablation Study on Replica [59]. The first 4 rows use DBA [67] instead of DSPO. Deforming the neural point cloud is important for reconstruction and rendering, the monocular prior helps with completeness and the DSPO layer improves both rendering and reconstruction metrics. All results are averaged over 8 scenes.

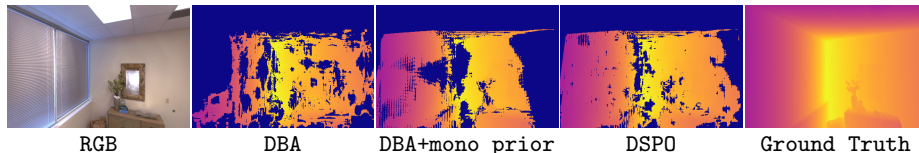


Fig. 5: Comparison of Estimated Depth. We show the depth output $\tilde{D}$ from the tracker. The pixels which are invalid (high error) are colored dark blue. DBA is the method that Droid-SLAM [67] uses. The DBA+mono prior strategy is used in HI-SLAM [84], *i.e.* the mono prior supervises all pixels directly. It is clear that our formulation (DSP0) provides the most consistent keyframe depth.

	Point-SLAM [55]	GO-SLAM [85]	SplaTAM* [26]	Ours
GPU Mem [GiB]	7.11	18.50	18.54	15.22
Avg. FPS	0.23	8.36	0.14	0.23

Table 7: Memory and Running Time Evaluation on Replica [59] room0. Our peak memory usage and runtime is comparable to existing works. We take the numbers from [69] except for ours. All methods except for ours are evaluated in RGB-D mode. All methods are evaluated using an NVIDIA RTX 3090 GPU.

Point-SLAM [55]. Although we both use the neural point cloud representation, we need to add more neural points to the scene to resolve the noisy depth, which Point-SLAM can avoid by having access to ground truth depth. Regarding runtime, we are comparable with Point-SLAM and SplaTAM. GO-SLAM has the fastest runtime, but as shown in Tab. 1 and Tab. 3, it sacrifices rendering and reconstruction quality for speed.

Limitations. We currently do not prune neural points that may be redundant or wrongly positioned as a result of depth noise. Most points are, however, corrected by the point cloud deformations, but point pruning would nonetheless be useful. Another limitation is that our construction of the final proxy depth D_c is rather simple and does not fuse the monocular and keyframe depth maps in an informed manner, by utilizing *e.g.* normal consistency. Finally, as future work, it is interesting to study how and when the frame-to-frame tracking paradigm can be boosted by frame-to-model techniques.

5 Conclusion

We proposed GIORIE-SLAM, a dense RGB-only SLAM system which utilizes a deformable neural point cloud for mapping and a globally optimized frame-to-frame tracking technique driven by optical flow. Importantly, the inclusion of a monocular depth prior into the tracking loop, to refine the scale and to regularize the erroneous keyframe depth predictions, leads to better rendering and mapping. By using the monocular prior for depth completion, mapping is further improved. Our experiments demonstrate that GIORIE-SLAM outperforms existing solutions regarding reconstruction and rendering accuracy while being on par or better with respect to tracking as well as runtime and memory usage.

Acknowledgments. We thank Wei Zhang for fruitful discussions.

GLORIE-SLAM: Globally Optimized RGB-only Implicit Encoding Point Cloud SLAM

Supplementary Material

Ganlin Zhang^{1*}, Erik Sandström^{1*}, Youmin Zhang^{2,3}, Manthan Patel¹,
Luc Van Gool^{1,4,5}, and Martin R. Oswald^{1,6}

¹ ETH Zürich, ² University of Bologna, ³ Rock Universe, ⁴ KU Leuven,
⁵ INSAIT, ⁶ University of Amsterdam

A Method

A.1 Mathematical Derivation of the DSPO Layer

In the main paper, we introduce the DSPO layer, which optimizes the disparity, scale and pose together. Here, we provide the detailed derivation of how to optimize the proposed DSPO objective using the Gauss-Newton algorithm with the Schur Complement.

As described in the main paper, we optimize the following two objective functions alternately. The first one (same as the DBA layer in [67]) optimizes the poses and disparity maps of the involved keyframes,

$$\arg \min_{\omega, d} \sum_{(i,j) \in \mathcal{E}} \left\| \tilde{p}_{ij} - K\omega_j^{-1}(\omega_i(1/d_i)K^{-1}[p_i, 1]^T) \right\|_{\Sigma_{ij}}^2. \quad (16)$$

The second objective optimizes the scale and shift factors of the mono prior and also the disparity maps of the keyframes,

$$\begin{aligned} \arg \min_{d^h, \theta, \gamma} \sum_{(i,j) \in \mathcal{E}} \left\| \tilde{p}_{ij} - K\omega_j^{-1}(\omega_i(1/d_i^h)K^{-1}[p_i, 1]^T) \right\|_{\Sigma_{ij}}^2 \\ + \alpha_1 \sum_{i \in \mathcal{V}} \left\| d_i^h - (\theta_i(1/D_i^{\text{mono}}) + \gamma_i) \right\|^2 + \alpha_2 \sum_{i \in \mathcal{V}} \left\| d_i^l - (\theta_i(1/D_i^{\text{mono}}) + \gamma_i) \right\|^2. \end{aligned} \quad (17)$$

For better readability, we do not show the conversion of 3D points to homogeneous coordinates in all equations.

DBA Optimization. First, we introduce how to solve the DBA optimization. To keep the notation consistent with [68], we do a mapping of the symbols we use in Eq. (16) such that

$$\{(p_i, \tilde{p}_{i,j}) \mid (i,j) \in \mathcal{E}\} \rightarrow \mathcal{M} = \{(x_i^k, x_j^k)\}_{k=1}^M, \quad (18)$$

where (x_i^k, x_j^k) is a corresponding pixel pair (a feature match found by the optical flow estimation), k is the index in the total set of matches $\mathcal{M}$. For each k there

* The two authors contributed equally to this paper.

are specific indices i and j indicating the corresponding keyframes of x_i^k and x_j^k , *i.e.* i and j are functions of k , but we write them as a subscript directly for simplicity. Furthermore, we map the camera poses

$$\{\omega_i\}_{i=1}^N \rightarrow \mathbf{T} = \{T_i \in SE(3) \mid 1 \leq i \leq N\} , \quad (19)$$

where $\mathbf{T}$ are the camera to world poses of the corresponding keyframes. Then we can rewrite the objective function in Eq. (16) as

$$f(\mathbf{T}, \mathbf{d}) = \frac{1}{2} \sum_{(x_i^k, x_j^k) \in \mathcal{M}} \|r_k(T_i, T_j, x_i^k, x_j^k, d_i^k)\|_{w_k}^2 , \quad (20)$$

where $\mathbf{d} = \{d_i^k\}_{k=1}^M$ are the disparities (*i.e.* inverse depth), d_i^k is the disparity of pixel x_i^k , where i is still a function of k , but we write it as a subscript for simplicity. We denote $\|\cdot\|_{w_k}$ as the Mahalanobis distance with weighting matrix w_k . $w_k = \text{diag}(w_k^1, w_k^2)$ is a 2×2 diagonal matrix of the per-pixel and per-direction uncertainties of the optical flow prediction. We add $1/2$ to the objective function for mathematical convenience as it does not change the optimal solution. $r_k(\cdot) \in \mathbb{R}^2$ is the reprojection error function and it is defined as

$$r_k(T_i, T_j, x_i^k, x_j^k, d_i^k) = x_j^k - K T_j^{-1} T_i (1/d_i^k) K^{-1} [x_i^k, 1]^T . \quad (21)$$

The formulation in Eq. (21) is the per-pixel equivalent of the part inside the norm $\|\cdot\|$ in Eq. (16).

To minimize Eq. (20), we resort to the Gauss-Newton algorithm. We can define the model parameters by the column vector $\beta = [T_j, T_i, d_i^k]^T \in \mathbb{R}^{13}$. We parameterize rotations and translations with the lie algebra $\text{se}(3)$, so each camera pose is parameterized by 6 values and the $d_i^k \in \mathbb{R}$. Despite being an RGB-only system, optimization is done on $\text{se}(3)$ because we fix the poses of the two first keyframes after initialization. This prevents gauge freedom *i.e.* drift in the global scale of the scene during optimization and therefore, $\text{se}(3)$ optimization is sufficient. We can define the model as

$$g_k(\beta) = \begin{bmatrix} g_k^1(\beta) \\ g_k^2(\beta) \end{bmatrix} = K T_j^{-1} T_i (1/d_i^k) K^{-1} [x_i^k, 1]^T . \quad (22)$$

since $g_k \in \mathbb{R}^2$ is not linear with respect to the parameters β we approximate g_k around a current estimate β_0 (at time 0) of the parameters with a first-order Taylor expansion. For each dimension of g_k , we get (here for the first dimension)

$$g_k^1(\beta) \approx g_k^1(\beta_0) + J_k(\beta - \beta_0) \quad (23)$$

where $J_k \in \mathbb{R}^{13}$ is the gradient row vector of g_k with respect to β at β_0 . Plugging back Eq. (23) into Eq. (21), we get (for the first dimension)

$$r_k^1(\beta) \approx \{x_j^k\}^1 - g_k^1(\beta_0) - J_k(\beta - \beta_0) = r_k^1(\beta_0) - J_k(\beta - \beta_0) , \quad (24)$$

where $\{x_j^k\}^1$ denotes the first dimension of x_j^k . The goal of the optimization is to minimize the squared residuals, *i.e.* for pixel k and dimension 1, the term $(r_k^1)^2$. We differentiate $(r_k^1)^2$ with respect to β and set the derivative to zero. This yields the expression

$$J_k^T J_k (\beta - \beta_0) + J_k^T r_k^1(\beta_0) = 0. \quad (25)$$

Now that we have derived the equation to solve (for β) for a single pixel and dimension, we can combine all observations from all pixels (further details can be found in section 2.4 in [68]). We achieve this by stacking all residuals as a column vector $\mathbf{r} = (r_1^1, r_1^2, \dots, r_M^1, r_M^2)^T \in \mathbb{R}^{2M}$. This yields the equivalent, vector-valued equation

$$\begin{aligned} \mathbf{J}^T \mathbf{W} \mathbf{J} (\beta - \beta_0) + \mathbf{J}^T \mathbf{W} \mathbf{r} &= \mathbf{J}^T \mathbf{W} \mathbf{J} \Delta \beta + \mathbf{J}^T \mathbf{W} \mathbf{r} \\ &= \mathbf{J}^T \mathbf{W} \mathbf{J} \begin{bmatrix} \Delta \mathbf{T} \\ \Delta \mathbf{d} \end{bmatrix} + \mathbf{J}^T \mathbf{W} \mathbf{r} = 0. \end{aligned} \quad (26)$$

Here, we combine the uncertainties into a diagonal matrix as

$$\mathbf{W} = \text{diag}(w_1^1, w_1^2, \dots, w_M^1, w_M^2). \quad (27)$$

The full Jacobian $\mathbf{J}$ has the shape $[2M \times (6N + M)]$, $\Delta \mathbf{T}$ is $[6N \times 1]$ and $\Delta \mathbf{d}$ is $[M \times 1]$. To solve Eq. (26), we rewrite it as

$$\mathbf{J}^T \mathbf{W} \mathbf{J} \begin{bmatrix} \Delta \mathbf{T} \\ \Delta \mathbf{d} \end{bmatrix} = -\mathbf{J}^T \mathbf{W} \mathbf{r}. \quad (28)$$

We now arrange the full Jacobian matrix $\mathbf{J}$ into two blocks as

$$\mathbf{J} = [\mathbf{J}_{\mathbf{T}} \ \mathbf{J}_{\mathbf{d}}], \quad (29)$$

where $\mathbf{J}_{\mathbf{T}} \in \mathbb{R}^{2M \times 6N}$ is the Jacobian block of $\mathbf{r}$ with respect to poses $\mathbf{T}$, and $\mathbf{J}_{\mathbf{d}} \in \mathbb{R}^{2M \times M}$ is the Jacobian block of $\mathbf{r}$ with respect to disparities $\mathbf{d}$. For the detailed derivation of $\mathbf{J}_{\mathbf{T}}$ and $\mathbf{J}_{\mathbf{d}}$, we refer to section 2.3.1 in [68]. This yields

$$\begin{bmatrix} \mathbf{J}_{\mathbf{T}}^T \mathbf{W} \mathbf{J}_{\mathbf{T}} & \mathbf{J}_{\mathbf{T}}^T \mathbf{W} \mathbf{J}_{\mathbf{d}} \\ \mathbf{J}_{\mathbf{d}}^T \mathbf{W} \mathbf{J}_{\mathbf{T}} & \mathbf{J}_{\mathbf{d}}^T \mathbf{W} \mathbf{J}_{\mathbf{d}} \end{bmatrix} \begin{bmatrix} \Delta \mathbf{T} \\ \Delta \mathbf{d} \end{bmatrix} = -\mathbf{J}^T \mathbf{W} \mathbf{r}. \quad (30)$$

We rewrite Eq. (30) to the following form,

$$\begin{bmatrix} \mathbf{B} & \mathbf{E} \\ \mathbf{E}^T & \mathbf{C} \end{bmatrix} \begin{bmatrix} \Delta \mathbf{T} \\ \Delta \mathbf{d} \end{bmatrix} = \begin{bmatrix} \mathbf{v} \\ \mathbf{w} \end{bmatrix}, \quad (31)$$

and use the Schur Complement to solve for $\Delta \mathbf{T}$ and $\Delta \mathbf{d}$,

$$\begin{aligned} \Delta \mathbf{T} &= [\mathbf{B} - \mathbf{E} \mathbf{C}^{-1} \mathbf{E}^T]^{-1} (\mathbf{v} - \mathbf{E} \mathbf{C}^{-1} \mathbf{w}) \\ \Delta \mathbf{d} &= \mathbf{C}^{-1} (\mathbf{w} - \mathbf{E}^T \Delta \mathbf{T}). \end{aligned} \quad (32)$$

Note that $\mathbf{C}$ here is diagonal since each d_i^k is only involved in r_k . Though $\mathbf{J}_{\mathbf{d}}$ is not perfectly diagonal (it has two non-zero values along each column), $\mathbf{J}_{\mathbf{d}}^T \mathbf{W} \mathbf{J}_{\mathbf{d}} = \mathbf{C} \in \mathbb{R}^{M \times M}$ is diagonal. Therefore, the inverse of $\mathbf{C}$ is easy to compute. Also, the size of $\mathbf{B} \in \mathbb{R}^{6N \times 6N}$ is relatively small and therefore it is tractable to compute $[\mathbf{B} - \mathbf{E} \mathbf{C}^{-1} \mathbf{E}^T]^{-1}$. This is achieved by Cholesky decomposition. Finally we can use $\Delta \mathbf{T}$ and $\Delta \mathbf{d}$ to update the poses and disparities accordingly.

Scale, Shift and Disparity Optimization. We approach the problem of joint scale, shift and disparity optimization in a similar manner as above. First, we rewrite the objective function in Eq. (17) as,

$$\begin{aligned} f(\mathbf{s}, \mathbf{d}_h) = & \frac{1}{2} \sum_{\substack{(x_i^k, x_j^k) \in \mathcal{M}_h \\ d_i^k \in \mathbf{d}_h}} \|r_k(T_i, T_j, x_i^k, x_j^k, d_i^k)\|_{w_k}^2 + \alpha_1 \|t_k(d_i^k, s_i)\|^2 \\ & + \frac{1}{2} \sum_{d_i^k \in \mathbf{d}_l} \alpha_2 \|t_k(d_i^k, s_i)\|^2 \end{aligned} \quad (33)$$

where $s_i = (\theta_i, \gamma_i)$ is the scale and shift of frame i . We denote the set

$$\mathcal{M}_h = \{(x_i^k, x_j^k)\}_{k=1}^H \quad (34)$$

as the set of pixels where d_i^k is deemed to have a high error as defined by the multi-view filter in Eq. (6). The shape of $\mathbf{d}_h$ is $[1 \times H]$ ($H \leq M$) and $\mathbf{d}_l$ denotes the set of disparities with a low error. Adding the cardinalities of $\mathbf{d}_h$ and $\mathbf{d}_l$ yields M *i.e.* $|\mathbf{d}_h| + |\mathbf{d}_l| = H + |\mathbf{d}_l| = M$. We denote the set of scales and shifts for all frames involves as $\mathbf{s} = (s_1, \dots, s_N)$. $t_k(\cdot) \in \mathbb{R}$ is the residual term for the regularization by the monocular depth prior,

$$t_k(d_i^k, s_i) = d_i^k - (\theta_i \cdot d_{\text{mono},i}^k + \gamma_i) \quad (35)$$

Now, we collect all the residuals as a column vector

$$\hat{\mathbf{r}} = (\dots, r_k^1, r_k^2, \dots, r_H^1, r_H^2, t_1, \dots, t_M)^T, \quad (36)$$

i.e. a collection of all r_k where k needs to ensure d_i^k is invalid (high error) and also all t . To define the corresponding linear system as in Eq. (30), we begin by defining the full Jacobian matrix $\hat{\mathbf{J}}$ as a collection of two blocks, similar to Eq. (29) *i.e.* $\hat{\mathbf{J}} = [\hat{\mathbf{J}}_{\mathbf{s}} \hat{\mathbf{J}}_{\mathbf{d}}]$, where $\hat{\mathbf{J}}_{\mathbf{s}} \in \mathbb{R}^{(2H+M) \times 2N}$ is the Jacobian block of $\hat{\mathbf{r}}$ with respect to the scales and shifts $\mathbf{s}$, and $\hat{\mathbf{J}}_{\mathbf{d}} \in \mathbb{R}^{(2H+M) \times H}$ is the Jacobian block of $\hat{\mathbf{r}}$ with respect to the invalid disparities $\mathbf{d}_h$. Here both $\hat{\mathbf{J}}_{\mathbf{s}}$ and $\hat{\mathbf{J}}_{\mathbf{d}}$ consist of two parts. The first part comes from $r_k(\cdot)$ and second part from $t_k(\cdot)$ *i.e.*,

$$\hat{\mathbf{J}}_{\mathbf{s}} = \begin{bmatrix} \hat{\mathbf{J}}_{\mathbf{s}}^r \\ \hat{\mathbf{J}}_{\mathbf{s}}^t \end{bmatrix} \quad \hat{\mathbf{J}}_{\mathbf{d}} = \begin{bmatrix} \hat{\mathbf{J}}_{\mathbf{d}}^r \\ \hat{\mathbf{J}}_{\mathbf{d}}^t \end{bmatrix} \quad \hat{\mathbf{J}} = \begin{bmatrix} \hat{\mathbf{J}}_{\mathbf{s}}^r & \hat{\mathbf{J}}_{\mathbf{d}}^r \\ \hat{\mathbf{J}}_{\mathbf{s}}^t & \hat{\mathbf{J}}_{\mathbf{d}}^t \end{bmatrix}. \quad (37)$$

$\hat{\mathbf{J}}_{\mathbf{d}}^r$ is the same as $\mathbf{J}_{\mathbf{d}}$ in Eq. (29) except that now it only contains the invalid disparities $\mathbf{d}_h$ instead of $\mathbf{d}$, *i.e.* the shape of $\hat{\mathbf{J}}_{\mathbf{d}}^r$ is $[2H \times H]$ instead of $[2M \times M]$. Note that we need to use a factor 2 here because each residual r_k has 2 dimensions. $\hat{\mathbf{J}}_{\mathbf{s}}^r \in \mathbb{R}^{2H \times 2N}$ is 0 since none of the s_i are involved in any of the $r_k(\cdot)$ residuals. The derivatives of $t_k(\cdot)$ with respect to θ_i , γ_i and d_i^k are,

$$\frac{\partial t_k}{\partial \theta_i} = -d_{\text{mono},i}^k \quad \frac{\partial t_k}{\partial \gamma_i} = -1 \quad \frac{\partial t_k}{\partial d_i^k} = 1 \quad . \quad (38)$$

Thus the form of $\hat{\mathbf{J}}_{\mathbf{s}}^t \in \mathbb{R}^{M \times 2N}$ is as follows,

$$\hat{\mathbf{J}}_{\mathbf{s}}^t = \begin{bmatrix} J_{1,1} & \dots & J_{1,N} \\ \vdots & \ddots & \vdots \\ J_{M,1} & \dots & J_{M,N} \end{bmatrix} \quad J_{k,i} = \begin{bmatrix} -d_{\text{mono},i}^k & -1 \end{bmatrix} \quad , \quad (39)$$

Finally, we construct $\hat{\mathbf{J}}_{\mathbf{d}}^t$ as follows. First define the diagonal square matrix $\mathbf{D} = \text{diag}(1, \dots, 1) \in \mathbb{R}^{M \times M}$. Then we define $\hat{\mathbf{J}}_{\mathbf{d}}^t$ as,

$$\hat{\mathbf{J}}_{\mathbf{d}}^t = [\dots D_k \dots] \in \mathbb{R}^{M \times H} \quad (40)$$

where D_k is the k_{th} column of $\mathbf{D}$, and k needs to ensure that $d_i^k \in \mathbf{d}_h$.

Next, we define the weighting matrix $\widehat{\mathbf{W}}$ as

$$\widehat{\mathbf{W}} = \text{diag}(\dots, w_k^1, w_k^2, \dots, w_H^1, w_H^2, \sigma_1, \dots, \sigma_M) \quad (41)$$

where k needs to ensure d_i^k is invalid (high error), and σ equals to α_1 or α_2 , depending on the corresponding disparity is invalid or valid (*i.e.* same α_1 and α_2 as used in Eq. (33)). Then, similar to Eq. (30), we use the Gauss-Newton method to form a linear equation,

$$\begin{bmatrix} \hat{\mathbf{J}}_{\mathbf{s}}^T \widehat{\mathbf{W}} \hat{\mathbf{J}}_{\mathbf{s}} & \hat{\mathbf{J}}_{\mathbf{s}}^T \widehat{\mathbf{W}} \hat{\mathbf{J}}_{\mathbf{d}} \\ \hat{\mathbf{J}}_{\mathbf{d}}^T \widehat{\mathbf{W}} \hat{\mathbf{J}}_{\mathbf{s}} & \hat{\mathbf{J}}_{\mathbf{d}}^T \widehat{\mathbf{W}} \hat{\mathbf{J}}_{\mathbf{d}} \end{bmatrix} \begin{bmatrix} \Delta \mathbf{s} \\ \Delta \mathbf{d} \end{bmatrix} = -\hat{\mathbf{J}}^T \widehat{\mathbf{W}} \hat{\mathbf{r}} \quad , \quad (42)$$

and again use the Schur Complement Eq. (31) to solve for $\Delta \mathbf{s}$ and $\Delta \mathbf{d}$. Note that the lower right corner block $\hat{\mathbf{J}}_{\mathbf{d}}^T \widehat{\mathbf{W}} \hat{\mathbf{J}}_{\mathbf{d}}$ is still diagonal. Therefore, it is easy to compute its inverse.

A.2 Interpolation of Neural Point Cloud

Basically the strategy we use for feature interpolation in the neural point cloud is the same as in Point-SLAM [55], but for the convenience of readers, we also

provide it here. In this section, we reuse the notation used in [55] for simplicity. Note that we redefine variables that may have been used before in this section.

As mentioned in the main paper, we use the same architecture for the occupancy decoder $\mathcal{H}$ and the color decoder $\mathcal{G}_\xi$ as [55] and use their provided pretrained and fixed middle geometric decoder $\mathcal{H}$. The decoder input is the 3D point x_i , to which we apply a learnable Gaussian positional encoding [64] to mitigate the limited band-width of MLPs, and the associated feature. We further denote $P^g(x_i)$ and $P^c(x_i)$ as the geometric and color features extracted at point x_i respectively. For each point p_i we use the corresponding per-pixel query radius $2r$, where r is computed according to Eq. (2). Within the radius $2r$, we require to find at least two neighbors. Otherwise, the point is given zero occupancy. We use the closest eight neighbors and use inverse squared distance weighting for the geometric features, *i.e.*

$$P^g(x_i) = \sum_k \frac{w_k}{\sum_k w_k} f_k^g \quad \text{with } w_k = \frac{1}{\|p_k - x_i\|^2} . \quad (43)$$

The color features are obtained similarly,

$$P^c(x_i) = \sum_k \frac{w_k}{\sum_k w_k} f_k^c \quad \text{with } w_k = \frac{1}{\|p_k - x_i\|^2} . \quad (44)$$

Note that we do not use the non-linear pre-processing on the extracted neighbor features as done in [55] since we found this to be unstable in the RGB-only setting. Next, we describe how the per-point occupancies s_i and colors t_i are used to render the per-pixel depth and color using volume rendering. We construct a weighting function, α_i as described in Eq. (45). This weight represents the discretized probability that the ray terminates at point x_i .

$$\alpha_i = s_{p_i} \prod_{j=1}^{i-1} (1 - s_{p_j}) . \quad (45)$$

The rendered depth is computed as the weighted average of the depth values along each ray, and equivalently for the color according to Eq. (46),

$$D_{\text{render}} = \sum_{i=1}^N \alpha_i z_i , \quad I_{\text{render}} = \sum_{i=1}^N \alpha_i t_i , \quad (46)$$

B More Experiments

To accompany the evaluations provided in the main paper, we provide further experiments in this section.

B.1 Full Evaluations Data

In Tab. 8, Tab. 9 and Tab. 10, we provide the full per scene results on all commonly reported metrics on Replica [59], TUM-RGBD [61] and ScanNet [10].

		Metric	R-0	R-1	R-2	0-0	0-1	0-2	0-3	0-4	Avg.
Reconstruction		Depth L1 ↓	2.98	2.52	3.30	3.29	2.76	3.72	4.16	3.11	3.24
		Accuracy ↓	2.84	3.07	3.05	2.98	2.06	3.32	3.34	2.92	2.96
		Completion ↓	4.65	3.55	3.64	2.39	3.43	4.54	4.57	4.78	3.95
		Comp. Rat. ↑	81.96	85.78	84.50	88.82	85.07	82.09	80.41	81.04	83.72
Render- ing	Key Frames	PSNR ↑	28.49	30.09	29.98	35.88	37.15	28.45	28.54	29.73	31.04
		SSIM ↑	0.96	0.97	0.96	0.98	0.99	0.97	0.97	0.97	0.97
		LPIPS ↓	0.13	0.13	0.14	0.09	0.08	0.15	0.11	0.15	0.12
	Every 5 Frames	PSNR ↑	27.15	28.85	28.56	33.95	36.27	26.78	26.95	28.11	29.58
		SSIM ↑	0.95	0.96	0.95	0.97	0.99	0.96	0.96	0.95	0.96
		LPIPS ↓	0.15	0.12	0.16	0.11	0.08	0.17	0.13	0.17	0.14
Track- ing	Key Frames Trajectory	ATE RMSE ↓	0.31	0.52	0.21	0.26	0.29	0.41	0.46	0.44	0.36
	Full Trajectory	ATE RMSE ↓	0.31	0.37	0.20	0.29	0.28	0.45	0.45	0.44	0.35

Table 8: Full Evaluation on Replica [59]. For Keyframes we only measure the performance on the frames which are used for mapping, which is the same strategy as existing methods. We also provide rendering results every 5 frames regardless whether they are keyframes or not, shown as **Every 5 Frames**. We show the ATE RMSE [cm] evaluation on the keyframes as well as on the full trajectory.

		Metric	f1/desk	f1/desk2	f1/room	f2/xyz	f3/office	Avg.
Render- ing	Key Frames	PSNR ↑	20.26	19.09	18.78	25.62	21.21	20.99
		SSIM ↑	0.87	0.82	0.79	0.96	0.84	0.85
		LPIPS ↓	0.31	0.38	0.38	0.09	0.32	0.30
	Every 5 Frames	PSNR ↑	17.16	15.93	16.17	21.72	17.25	17.65
		SSIM ↑	0.79	0.72	0.72	0.92	0.73	0.77
		LPIPS ↓	0.40	0.44	0.42	0.15	0.44	0.37
Track- ing	Key Frames Trajectory	ATE RMSE ↓	1.92	3.05	4.43	0.23	1.41	2.21
	Full Trajectory	ATE RMSE ↓	1.65	2.79	4.16	0.22	1.44	2.05

Table 9: Full Evaluation on TUM-RGBD [61].

		Metric	0000	0054	0059	0106	0169	0181	0207	0233	Avg.
Render- ing	Key Frames	PSNR↑	23.42	25.68	20.66	20.41	25.23	21.28	23.68	22.29	22.83
		SSIM ↑	0.87	0.87	0.83	0.84	0.91	0.76	0.85	0.84	0.84
		LPIPS ↓	0.26	0.30	0.31	0.31	0.21	0.44	0.29	0.30	0.30
	Every 5 Frames	PSNR↑	21.80	23.75	18.66	17.54	22.85	19.14	21.52	20.36	20.70
		SSIM ↑	0.83	0.83	0.77	0.77	0.88	0.71	0.81	0.78	0.80
		LPIPS ↓	0.30	0.34	0.36	0.42	0.25	0.47	0.32	0.36	0.35
Track- ing	Key Frames Trajectory	ATE RMSE ↓	5.66	9.17	9.48	7.03	8.72	8.42	7.47	4.97	7.61
	Full Trajectory	ATE RMSE ↓	5.57	9.50	9.11	7.09	8.26	8.39	7.53	5.17	7.58

Table 10: Full Evaluation on ScanNet [10].

The reconstruction results are only measured on Replica since the other two datasets are real world datasets which lack quality ground truth meshes.

	Metric	R-0	R-1	R-2	0-0	0-1	0-2	0-3	0-4	Avg.
Recon- struction	Depth L1 ↓	2.16	1.41	1.85	1.64	1.77	3.23	5.71	1.93	2.46
	Accuracy ↓	2.02	1.31	1.43	1.22	0.95	3.47	4.99	1.77	2.15
	Completion ↓	3.26	2.71	2.59	1.56	2.19	3.80	3.23	3.52	2.86
	Comp. Rat. ↑	88.27	89.98	89.65	94.06	91.11	85.17	85.19	85.65	88.64
Rendering	PSNR ↑	29.98	31.98	32.86	37.29	38.03	30.34	29.40	33.37	32.91
	SSIM ↑	0.97	0.98	0.98	0.99	0.99	0.98	0.98	0.98	0.98
	LPIPS ↓	0.10	0.11	0.09	0.08	0.06	0.13	0.10	0.12	0.10
Tracking	ATE RMSE ↓	0.38	0.36	0.24	0.24	0.45	0.37	0.42	0.44	0.36

Table 11: Full Evaluations on Replica [59] with ground truth depth. Both reconstruction and rendering results improve significantly with the ground truth depth, suggesting that our method is bounded by the quality of current day monocular depth estimation. Since we do not require any extra training or fine-tuning of the monocular depth estimator, it is easy to plug in a better estimator once available. Tracking performance does not change much.

For the rendering metrics, we not only report the results on the keyframes (the frames used for mapping), but also the every-5-frames results. The every-5-frames rendering performance is very close compared to the keyframe rendering results, which suggests that our method can generalize to unseen views.

We also show the trajectory accuracy measurement of both keyframes and all the full trajectory, which is obtained by first linear interpolation between keyframes and using optical flow to refine. The accuracy of these two trajectories are similar. In the main paper, the data we report is always measured on the full trajectory.

B.2 Influence of Monocular Depth Prior

While we show that the monocular depth prior improves the geometric estimation capability of our framework, it may still be erroneous. To better understand the accuracy of monocular depth prior, we replace the monocular depth prior with the ground truth sensor depth instead. This experiment acts as the upper bound of our method if the monocular depth prior is perfect. The experiments are done on Replica [59] and are shown in Tab. 11. Compared with the standard setting with the monocular prior, the ground truth depth setting gives improvements on both reconstruction and rendering quality, which reveals that our method still has potential to achieve better mapping results once a better monocular depth prior is available. Since our method does not require further training or fine-tuning for the monocular depth prior, it is quite easy to just replace the current off-the-shelf monocular depth estimator with a better one.

B.3 Scale and Shift in Depth Space or Disparity Space

In our paper, we align the scale and shift of the monocular depth prior in depth space with the valid estimated depth, as shown in Eq. (8) in the main paper. However, there are some other works [52] that suggest that the alignment should

	PSNR [dB] $\uparrow$	Depth L1 [cm] $\downarrow$	Accuracy [cm] $\downarrow$	Completion [cm] $\downarrow$	Completion Ratio [%] $\uparrow$
Disparity Space	27.96	2.53	3.15	3.61	85.61
Depth Space	30.09	2.52	3.07	3.55	85.78

Table 12: Comparison of Scale and Shift Alignment in Different Spaces on Replica [59] room1. There is marginal performance difference between disparity and depth alignment. Using depth alignment is marginally better in all metrics.

be done in disparity space. We therefore compare the difference between using disparity space or depth space alignment. The results are shown in Tab. 12. The results show that in general the space choice does not influence the results much, especially for geometry, but depth space still performs a bit better. Thus, we use depth space to do the scale and shift alignment in our proposed method.

B.4 Additional Qualitative Renderings

In Fig. 6 we show additional renderings from the ScanNet dataset where our method is compared to GO-SLAM [85].

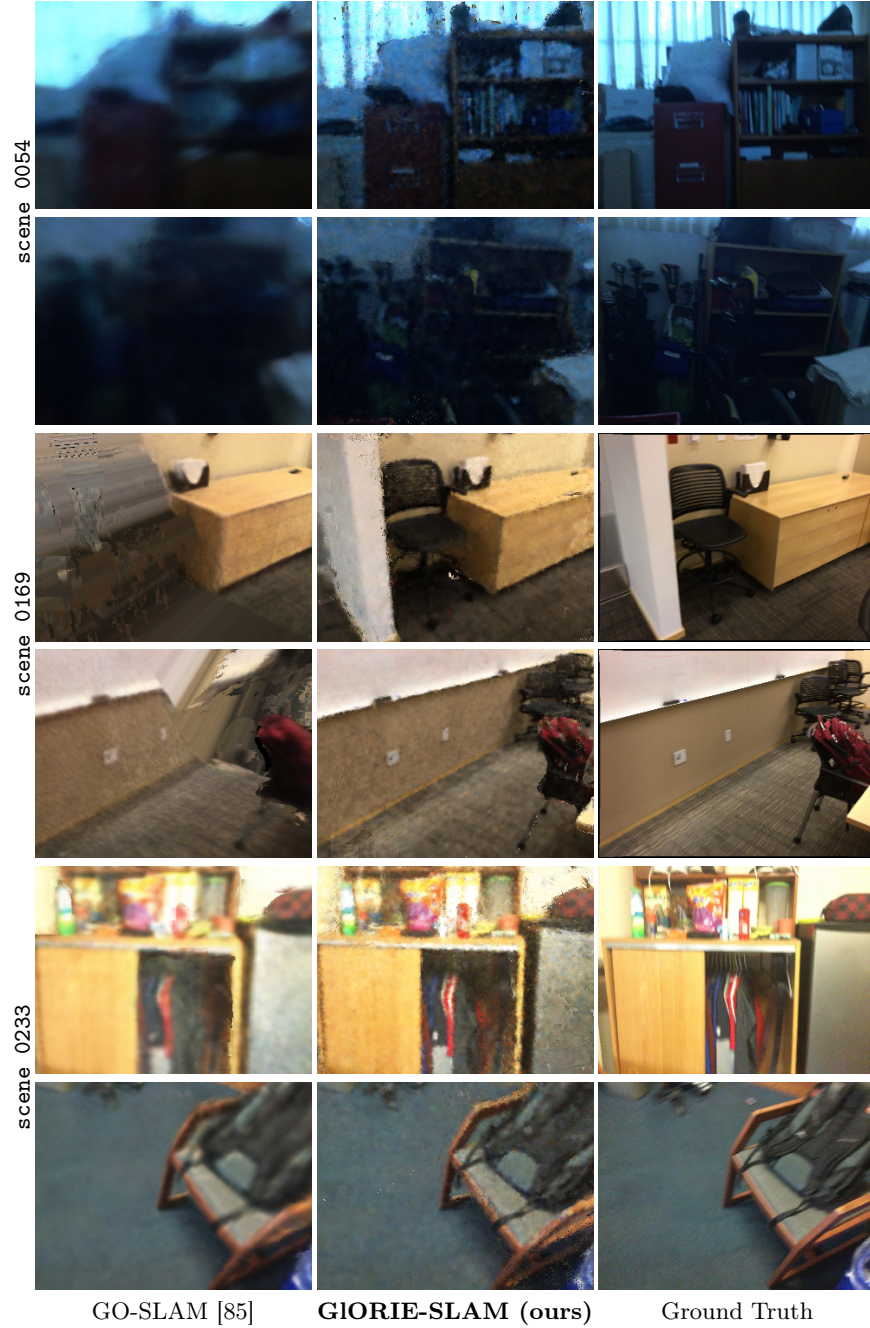


Fig. 6: Additional Rendering Results on ScanNet [10]. Thanks to the deformable neural point cloud, our method can maintain a globally consistent map and render high-frequency details with high fidelity.

References

1. Azinović, D., Martin-Brualla, R., Goldman, D.B., Nießner, M., Thies, J.: Neural rgb-d surface reconstruction. In: IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 6290–6301 (2022)
2. Bosse, M., Newman, P., Leonard, J., Soika, M., Feiten, W., Teller, S.: An atlas framework for scalable mapping. In: 2003 IEEE International Conference on Robotics and Automation (Cat. No. 03CH37422). vol. 2, pp. 1899–1906. IEEE (2003)
3. Božič, A., Palafox, P., Thies, J., Dai, A., Nießner, M.: Transformerfusion: Monocular rgb scene reconstruction using transformers. arXiv preprint arXiv:2107.02191 (2021)
4. Cao, Y.P., Kobbelt, L., Hu, S.M.: Real-time high-accuracy three-dimensional reconstruction with consumer rgb-d cameras. *ACM Transactions on Graphics (TOG)* **37**(5), 1–16 (2018)
5. Chen, J., Bautembach, D., Izadi, S.: Scalable real-time volumetric surface reconstruction. *ACM Transactions on Graphics (ToG)* **32**(4), 1–16 (2013)
6. Cho, H.M., Jo, H., Kim, E.: Sp-slam: Surfel-point simultaneous localization and mapping. *IEEE/ASME Transactions on Mechatronics* **27**(5), 2568–2579 (2021)
7. Choi, S., Zhou, Q.Y., Koltun, V.: Robust reconstruction of indoor scenes. In: IEEE Conference on Computer Vision and Pattern Recognition. pp. 5556–5565 (2015)
8. Chung, C.M., Tseng, Y.C., Hsu, Y.C., Shi, X.Q., Hua, Y.H., Yeh, J.F., Chen, W.C., Chen, Y.T., Hsu, W.H.: Orbeez-slam: A real-time monocular visual slam with orb features and nerf-realized mapping. arXiv preprint arXiv:2209.13274 (2022)
9. Curless, B., Levoy, M.: Volumetric method for building complex models from range images. In: SIGGRAPH Conference on Computer Graphics. ACM (1996)
10. Dai, A., Chang, A.X., Savva, M., Halber, M., Funkhouser, T., Nießner, M.: ScanNet: Richly-annotated 3D reconstructions of indoor scenes. In: Conference on Computer Vision and Pattern Recognition (CVPR). IEEE/CVF (2017). <https://doi.org/10.1109/CVPR.2017.261>, <http://arxiv.org/abs/1702.04405>
11. Dai, A., Nießner, M., Zollhöfer, M., Izadi, S., Theobalt, C.: Bundlefusion: Real-time globally consistent 3d reconstruction using on-the-fly surface reintegration. *ACM Transactions on Graphics (ToG)* **36**(4), 1 (2017)
12. Eftekhari, A., Sax, A., Malik, J., Zamir, A.: Omnidata: A scalable pipeline for making multi-task mid-level vision datasets from 3d scans. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 10786–10796 (2021)
13. Endres, F., Hess, J., Engelhard, N., Sturm, J., Cremers, D., Burgard, W.: An evaluation of the rgb-d slam system. In: 2012 IEEE international conference on robotics and automation. pp. 1691–1696. IEEE (2012)
14. Engel, J., Schöps, T., Cremers, D.: Lsd-slam: Large-scale direct monocular slam. In: European conference on computer vision. pp. 834–849. Springer (2014)
15. Fioraio, N., Taylor, J., Fitzgibbon, A., Di Stefano, L., Izadi, S.: Large-scale and drift-free surface reconstruction using online subvolume registration. In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition. pp. 4475–4483 (2015)
16. Henry, P., Fox, D., Bhowmik, A., Mongia, R.: Patch volumes: Segmentation-based consistent mapping with rgb-d cameras. In: 2013 International Conference on 3D Vision-3DV 2013. pp. 398–405. IEEE (2013)
17. Henry, P., Krainin, M., Herbst, E., Ren, X., Fox, D.: Rgb-d mapping: Using kinect-style depth cameras for dense 3d modeling of indoor environments. *The international journal of Robotics Research* **31**(5), 647–663 (2012)

18. Hu, J., Mao, M., Bao, H., Zhang, G., Cui, Z.: CP-SLAM: Collaborative neural point-based SLAM system. In: Thirty-seventh Conference on Neural Information Processing Systems (2023), <https://openreview.net/forum?id=dFSeZm6dTC>
19. Hu, P., Han, Z.: Learning neural implicit through volume rendering with attentive depth fusion priors. In: Advances in Neural Information Processing Systems (NeurIPS) (2023)
20. Hua, T., Bai, H., Cao, Z., Liu, M., Tao, D., Wang, L.: Hi-map: Hierarchical factorized radiance field for high-fidelity monocular dense mapping. arXiv preprint arXiv:2401.03203 (2024)
21. Hua, T., Bai, H., Cao, Z., Wang, L.: Fmapping: Factorized efficient neural field mapping for real-time dense rgb slam. arXiv preprint arXiv:2306.00579 (2023)
22. Huang, H., Li, L., Cheng, H., Yeung, S.K.: Photo-slam: Real-time simultaneous localization and photorealistic mapping for monocular, stereo, and rgb-d cameras. arXiv preprint arXiv:2311.16728 (2023)
23. Inmann, M., Zollhöfer, M., Nießner, M., Theobalt, C., Stamminger, M.: Volumedeform: Real-time volumetric non-rigid reconstruction. In: Computer Vision—ECCV 2016: 14th European Conference, Amsterdam, The Netherlands, October 11–14, 2016, Proceedings, Part VIII 14. pp. 362–379. Springer (2016)
24. Kähler, O., Prisacariu, V.A., Murray, D.W.: Real-time large-scale dense 3d reconstruction with loop closure. In: Computer Vision—ECCV 2016: 14th European Conference, Amsterdam, The Netherlands, October 11–14, 2016, Proceedings, Part VIII 14. pp. 500–516. Springer (2016)
25. Kähler, O., Prisacariu, V.A., Ren, C.Y., Sun, X., Torr, P.H.S., Murray, D.W.: Very high frame rate volumetric integration of depth images on mobile devices. IEEE Trans. Vis. Comput. Graph. **21**(11), 1241–1250 (2015). <https://doi.org/10.1109/TVCG.2015.2459891>, <https://doi.org/10.1109/TVCG.2015.2459891>
26. Keetha, N., Karhade, J., Jatavallabhula, K.M., Yang, G., Scherer, S., Ramanan, D., Luiten, J.: Splatam: Splat, track and map 3d gaussians for dense rgb-d slam. arXiv preprint (2023)
27. Keller, M., Lefloch, D., Lambers, M., Izadi, S., Weyrich, T., Kolb, A.: Real-time 3d reconstruction in dynamic scenes using point-based fusion. In: International Conference on 3D Vision (3DV). pp. 1–8. IEEE (2013)
28. Kerl, C., Sturm, J., Cremers, D.: Dense visual slam for rgb-d cameras. In: 2013 IEEE/RSJ International Conference on Intelligent Robots and Systems. pp. 2100–2106. IEEE (2013)
29. Li, H., Gu, X., Yuan, W., Yang, L., Dong, Z., Tan, P.: Dense rgb slam with neural implicit maps. In: Proceedings of the International Conference on Learning Representations (2023), <https://openreview.net/forum?id=QUK1Ex1bba>
30. Li, K., Tang, Y., Prisacariu, V.A., Torr, P.H.: Bnv-fusion: Dense 3d reconstruction using bi-level neural volume fusion. In: IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 6166–6175 (2022)
31. Li, K., Niemeyer, M., Navab, N., Tombari, F.: Dns slam: Dense neural semantic-informed slam. arXiv preprint arXiv:2312.00204 (2023)
32. Li, M., He, J., Wang, Y., Wang, H.: End-to-end rgb-d slam with multi-mlps dense neural implicit representations. IEEE Robotics and Automation Letters (2023)
33. Liso, L., Sandström, E., Yugay, V., Van Gool, L., Oswald, M.R.: Loopy-slam: Dense neural slam with loop closures. arXiv preprint arXiv:2402.09944 (2024)
34. Liu, L., Gu, J., Zaw Lin, K., Chua, T.S., Theobalt, C.: Neural sparse voxel fields. Advances in Neural Information Processing Systems **33**, 15651–15663 (2020)
35. Lorensen, W.E., Cline, H.E.: Marching cubes: A high resolution 3d surface construction algorithm. ACM siggraph computer graphics **21**(4), 163–169 (1987)

36. Mahdi Johari, M., Carta, C., Fleuret, F.: Eslam: Efficient dense slam system based on hybrid representation of signed distance fields. *arXiv e-prints* pp. arXiv-2211 (2022)
37. Maier, R., Schaller, R., Cremers, D.: Efficient online surface correction for real-time large-scale 3d reconstruction. *arxiv* 2017. *arXiv preprint* arXiv:1709.03763 (2017)
38. Maier, R., Sturm, J., Cremers, D.: Submap-based bundle adjustment for 3d reconstruction from rgb-d data. In: *Pattern Recognition: 36th German Conference, GCPR 2014, Münster, Germany, September 2-5, 2014, Proceedings* 36. pp. 54–65. Springer (2014)
39. Mao, Y., Yu, X., Wang, K., Wang, Y., Xiong, R., Liao, Y.: Ngel-slam: Neural implicit representation-based global consistent low-latency slam system. *arXiv preprint* arXiv:2311.09525 (2023)
40. Marniok, N., Johannsen, O., Goldluecke, B.: An efficient octree design for local variational range image fusion. In: *German Conference on Pattern Recognition (GCPR)*. pp. 401–412. Springer (2017)
41. Matsuki, H., Murai, R., Kelly, P.H., Davison, A.J.: Gaussian splatting slam. *arXiv preprint* arXiv:2312.06741 (2023)
42. Matsuki, H., Sucar, E., Laidow, T., Wada, K., Scona, R., Davison, A.J.: imode: Real-time incremental monocular dense mapping using neural field. In: *2023 IEEE International Conference on Robotics and Automation (ICRA)*. pp. 4171–4177. IEEE (2023)
43. Matsuki, H., Tatenno, K., Niemeyer, M., Tombari, F.: Newton: Neural view-centric mapping for on-the-fly large-scale slam. *arXiv preprint* arXiv:2303.13654 (2023)
44. Müller, T., Evans, A., Schied, C., Keller, A.: Instant neural graphics primitives with a multiresolution hash encoding. *arXiv preprint* arXiv:2201.05989 (2022)
45. Mur-Artal, R., Tardós, J.D.: Orb-slam2: An open-source slam system for monocular, stereo, and rgb-d cameras. *IEEE Transactions on Robotics* **33**(5), 1255–1262 (2017)
46. Naumann, J., Xu, B., Leutenegger, S., Zuo, X.: Nerf-vo: Real-time sparse visual odometry with neural radiance fields. *arXiv preprint* arXiv:2312.13471 (2023)
47. Newcombe, R.A., Izadi, S., Hilliges, O., Molyneaux, D., Kim, D., Davison, A.J., Kohli, P., Shotton, J., Hodges, S., Fitzgibbon, A.W.: Kinectfusion: Real-time dense surface mapping and tracking. In: *ISMAR*. vol. 11, pp. 127–136 (2011)
48. Nießner, M., Zollhöfer, M., Izadi, S., Stamminger, M.: Real-time 3d reconstruction at scale using voxel hashing. *ACM Transactions on Graphics (TOG)* **32** (11 2013). <https://doi.org/10.1145/2508363.2508374>
49. Oleynikova, H., Taylor, Z., Fehr, M., Siegwart, R., Nieto, J.I.: Voxblox: Incremental 3d euclidean signed distance fields for on-board MAV planning. In: *2017 IEEE/RSJ International Conference on Intelligent Robots and Systems, IROS 2017, Vancouver, BC, Canada, September 24-28, 2017*. pp. 1366–1373. IEEE (2017). <https://doi.org/10.1109/IROS.2017.8202315>, <https://doi.org/10.1109/IROS.2017.8202315>
50. Ortiz, J., Clegg, A., Dong, J., Sucar, E., Novotny, D., Zollhoefer, M., Mukadam, M.: isdf: Real-time neural signed distance fields for robot perception. *arXiv preprint* arXiv:2204.02296 (2022)
51. Peng, S., Niemeyer, M., Mescheder, L., Pollefeys, M., Geiger, A.: Convolutional Occupancy Networks. In: *European Conference Computer Vision (ECCV)*. CVM (2020), <https://www.microsoft.com/en-us/research/publication/convolutional-occupancy-networks/>

52. Ranftl, R., Lasinger, K., Hafner, D., Schindler, K., Koltun, V.: Towards robust monocular depth estimation: Mixing datasets for zero-shot cross-dataset transfer. *IEEE transactions on pattern analysis and machine intelligence* **44**(3), 1623–1637 (2020)
53. Reijgwart, V., Millane, A., Oleynikova, H., Siegwart, R., Cadena, C., Nieto, J.: Voxgraph: Globally consistent, volumetric mapping using signed distance function submaps. *IEEE Robotics and Automation Letters* **5**(1), 227–234 (2019)
54. Rosinol, A., Leonard, J.J., Carlone, L.: NeRF-SLAM: Real-Time Dense Monocular SLAM with Neural Radiance Fields. *arXiv* (2022), <http://arxiv.org/abs/2210.13641>
55. Sandström, E., Li, Y., Van Gool, L., Oswald, M.R.: Point-slam: Dense neural point cloud-based slam. In: *International Conference on Computer Vision (ICCV)*. IEEE/CVF (2023)
56. Sandström, E., Ta, K., Gool, L.V., Oswald, M.R.: Uncle-SLAM: Uncertainty learning for dense neural SLAM. In: *International Conference on Computer Vision Workshops (ICCVW)* (2023)
57. Schops, T., Sattler, T., Pollefeys, M.: BAD SLAM: Bundle adjusted direct RGB-D SLAM. In: *CVF/IEEE Conference on Computer Vision and Pattern Recognition (CVPR)* (2019)
58. Steinbrucker, F., Kerl, C., Cremers, D.: Large-scale multi-resolution surface reconstruction from rgb-d sequences. In: *IEEE International Conference on Computer Vision*. pp. 3264–3271 (2013)
59. Straub, J., Whelan, T., Ma, L., Chen, Y., Wijmans, E., Green, S., Engel, J.J., Mur-Artal, R., Ren, C., Verma, S., et al.: The replica dataset: A digital replica of indoor spaces. *arXiv preprint arXiv:1906.05797* (2019)
60. Stückler, J., Behnke, S.: Multi-resolution surfel maps for efficient dense 3d modeling and tracking. *Journal of Visual Communication and Image Representation* **25**(1), 137–147 (2014)
61. Sturm, J., Engelhard, N., Endres, F., Burgard, W., Cremers, D.: A benchmark for the evaluation of RGB-D SLAM systems. In: *International Conference on Intelligent Robots and Systems (IROS)*. IEEE/RSJ (2012). <https://doi.org/10.1109/IROS.2012.6385773>, <http://ieeexplore.ieee.org/document/6385773/>
62. Sucar, E., Liu, S., Ortiz, J., Davison, A.J.: iMAP: Implicit Mapping and Positioning in Real-Time. In: *International Conference on Computer Vision (ICCV)*. IEEE/CVF (2021). <https://doi.org/10.1109/ICCV48922.2021.00617>, <https://ieeexplore.ieee.org/document/9710431/>
63. Sun, J., Xie, Y., Chen, L., Zhou, X., Bao, H.: Neuralrecon: Real-time coherent 3d reconstruction from monocular video. In: *IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 15598–15607 (2021)
64. Tancik, M., Srinivasan, P., Mildenhall, B., Fridovich-Keil, S., Raghavan, N., Singhal, U., Ramamoorthi, R., Barron, J., Ng, R.: Fourier features let networks learn high frequency functions in low dimensional domains. *Advances in Neural Information Processing Systems* **33**, 7537–7547 (2020)
65. Tang, Y., Zhang, J., Yu, Z., Wang, H., Xu, K.: Mips-fusion: Multi-implicit-submaps for scalable and robust online neural rgb-d reconstruction. *arXiv preprint arXiv:2308.08741* (2023)
66. Teed, Z., Deng, J.: Raft: Recurrent all-pairs field transforms for optical flow. In: *Computer Vision–ECCV 2020: 16th European Conference, Glasgow, UK, August 23–28, 2020, Proceedings, Part II* 16. pp. 402–419. Springer (2020)

67. Teed, Z., Deng, J.: Droid-slam: Deep visual slam for monocular, stereo, and rgb-d cameras. *Advances in neural information processing systems* **34**, 16558–16569 (2021)
68. Teed, Z., et al.: Optimization Inspired Neural Networks for Multiview 3D Reconstruction. Ph.D. thesis, Princeton University (2022)
69. Tosi, F., Zhang, Y., Gong, Z., Sandström, E., Mattoccia, S., Oswald, M.R., Poggi, M.: How nerfs and 3d gaussian splatting are reshaping slam: a survey (2024)
70. Wang, H., Wang, J., Liang, W.: Online reconstruction of indoor scenes from rgb-d streams. In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*. pp. 3271–3279 (2016)
71. Wang, H., Wang, J., Agapito, L.: Co-slam: Joint coordinate and sparse parametric encodings for neural real-time slam. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. pp. 13293–13302 (June 2023)
72. Wang, Z., Bovik, A.C., Sheikh, H.R., Simoncelli, E.P.: Image quality assessment: from error visibility to structural similarity. *IEEE transactions on image processing* **13**(4), 600–612 (2004)
73. Weder, S., Schonberger, J., Pollefeys, M., Oswald, M.R.: Routedfusion: Learning real-time depth map fusion. In: *IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 4887–4897 (2020)
74. Weder, S., Schonberger, J.L., Pollefeys, M., Oswald, M.R.: Neuralfusion: Online depth fusion in latent space. In: *IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 3162–3172 (2021)
75. Whelan, T., Leutenegger, S., Salas-Moreno, R., Glocker, B., Davison, A.: Elasticfusion: Dense slam without a pose graph. In: *Robotics: Science and Systems (RSS)* (2015)
76. Xinyang, L., Yijin, L., Yanbin, T., Hujun, B., Guofeng, Z., Yinda, Z., Zhaopeng, C.: Multi-modal neural radiance field for monocular dense slam with a light-weight tof sensor. In: *International Conference on Computer Vision (ICCV)* (2023)
77. Yan, C., Qu, D., Wang, D., Xu, D., Wang, Z., Zhao, B., Li, X.: Gs-slam: Dense visual slam with 3d gaussian splatting. *arXiv preprint arXiv:2311.11700* (2023)
78. Yan, Z., Ye, M., Ren, L.: Dense visual slam with probabilistic surfel map. *IEEE transactions on visualization and computer graphics* **23**(11), 2389–2398 (2017)
79. Yang, X., Li, H., Zhai, H., Ming, Y., Liu, Y., Zhang, G.: Vox-fusion: Dense tracking and mapping with voxel-based neural implicit representation. In: *IEEE International Symposium on Mixed and Augmented Reality (ISMAR)*. pp. 499–507. IEEE (2022)
80. Yang, X., Ming, Y., Cui, Z., Calway, A.: Fd-slam: 3-d reconstruction using features and dense matching. In: *2022 International Conference on Robotics and Automation (ICRA)*. pp. 8040–8046. IEEE (2022)
81. Yugay, V., Li, Y., Gevers, T., Oswald, M.R.: Gaussian-slam: Photo-realistic dense slam with gaussian splatting (2023)
82. Zhang, H., Chen, G., Wang, Z., Wang, Z., Sun, L.: Dense 3d mapping for indoor environment based on feature-point slam method. In: *2020 the 4th International Conference on Innovation in Artificial Intelligence*. pp. 42–46 (2020)
83. Zhang, R., Isola, P., Efros, A.A., Shechtman, E., Wang, O.: The unreasonable effectiveness of deep features as a perceptual metric. In: *IEEE conference on computer vision and pattern recognition*. pp. 586–595 (2018)
84. Zhang, W., Sun, T., Wang, S., Cheng, Q., Haala, N.: Hi-slam: Monocular real-time dense mapping with hybrid implicit fields. *IEEE Robotics and Automation Letters* (2023)

85. Zhang, Y., Tosi, F., Mattoccia, S., Poggi, M.: Go-slam: Global optimization for consistent 3d instant reconstruction. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 3727–3737 (2023)
86. Zhou, H., Guo, Z., Liu, S., Zhang, L., Wang, Q., Ren, Y., Li, M.: Mod-slam: Monocular dense mapping for unbounded 3d scene reconstruction (2024)
87. Zhu, S., Wang, G., Blum, H., Liu, J., Song, L., Pollefeys, M., Wang, H.: Sni-slam: Semantic neural implicit slam. arXiv preprint arXiv:2311.11016 (2023)
88. Zhu, Z., Peng, S., Larsson, V., Cui, Z., Oswald, M.R., Geiger, A., Pollefeys, M.: Nicer-slam: Neural implicit scene encoding for rgb slam. arXiv preprint arXiv:2302.03594 (2023)
89. Zhu, Z., Peng, S., Larsson, V., Xu, W., Bao, H., Cui, Z., Oswald, M.R., Pollefeys, M.: Nice-slam: Neural implicit scalable encoding for slam. In: IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 12786–12796 (2022)
90. Zou, Z.X., Huang, S.S., Cao, Y.P., Mu, T.J., Shan, Y., Fu, H.: Mononeuralfusion: Online monocular neural 3d reconstruction with geometric priors. arXiv preprint arXiv:2209.15153 (2022)