

Revealing Occlusions with 4D Neural Fields

Basile Van Hoorick¹ Purva Tendulkar¹ Dídac Surís¹ Dennis Park² Simon Stent² Carl Vondrick¹

¹Columbia University ²Toyota Research Institute

{basile, purvaten, didacsuris, vondrick}@cs.columbia.edu, {dennis.park, simon.stent}@tri.global

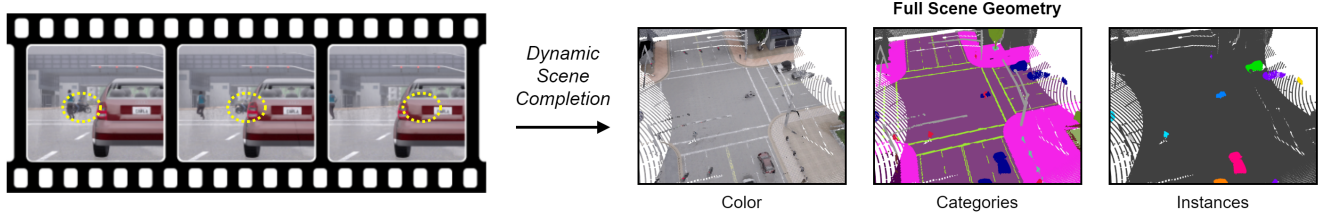


Figure 1. **Video Occlusions** – Although the motorcycle (circled in yellow) becomes fully occluded in the video, we can still perform many visual recognition tasks, such as predicting its location, reconstructing its appearance, and classifying its semantic category. This paper introduces a video representation architecture that is able to learn to perform all of these occlusion reasoning tasks. We show example inputs and ground truths of the proposed dynamic scene completion framework.

Abstract

For computer vision systems to operate in dynamic situations, they need to be able to represent and reason about object permanence. We introduce a framework for learning to estimate 4D visual representations from monocular RGB-D video, which is able to persist objects, even once they become obstructed by occlusions. Unlike traditional video representations, we encode point clouds into a continuous representation, which permits the model to attend across the spatiotemporal context to resolve occlusions. On two large video datasets that we release along with this paper, our experiments show that the representation is able to successfully reveal occlusions for several tasks, without any architectural changes. Visualizations show that the attention mechanism automatically learns to follow occluded objects. Since our approach can be trained end-to-end and is easily adaptable, we believe it will be useful for handling occlusions in many video understanding tasks. Data, code, and models are available at occlusions.cs.columbia.edu.

1. Introduction

When an object becomes occluded in video, its location and visual structure is often still predictable. In several studies, developmental psychologists have been able to demonstrate that shortly after birth, children learn how objects persist during occlusions [2, 5, 42, 52], and evidence suggests

that animals perform similar reasoning too [32, 41].¹ For example, although the yellow orb in Figure 1 disappears behind other objects, its location, geometry, and appearance remain evident to you. Occlusions are fundamental to computer vision, and predicting the contents behind them underlies many applications in video analysis.

The field has developed a number of deep learning methods to operate on point clouds [19, 44, 66, 68, 69], which due to their attractive properties, have emerged as the representation of choice for numerous 3D tasks. Point clouds are sparse, making them particularly scalable to large scenes. However, to solve the problem in Figure 1, we need a video representation that (1) uses evidence from the previous frames in order to (2) generate the new points that are not observed in the subsequent frames. Since point clouds are possible to collect at scale [7, 9], we believe they are an excellent source of data for learning to predict behind occlusions in video. However, the representation must also have the capacity to *create* points conditioned on their context.

This paper introduces an architecture for learning to predict 4D point clouds from an RGB-D video camera. The key to our approach is a continuous neural field representation of a point cloud, which uses an attention mechanism to condition the full space on the observations. Since the representation is continuous, the approach can learn to produce points anywhere in spacetime, allowing for high-fidelity reconstructions of complex scenes. Where there are occlusions and missing observations, the representation is able

¹See “What The Fluff Challenge” on YouTube.

to use attention to find the object and missing scene structure when it was last visible, and subsequently put the right points in the right place.

Experiments show that our video representation learns to successfully perform many occlusion reasoning tasks, such as visual reconstruction, geometry estimation, tracking, and semantic segmentation. The same method works for these tasks without architectural changes. On two different datasets, we show the approach remains robust for both highly cluttered scenes and objects of various sizes. Though we train the representation without ground truth correspondence, visualizations show that the attention mechanism automatically learns to follow objects through occlusions.

There are three principal contributions in this paper. Firstly, we propose the new fundamental task of **4D dynamic scene completion**, which forms a basis for spatiotemporal reasoning tasks. Secondly, we present new benchmarks to evaluate scene completion and object permanence in cluttered situations. Thirdly, we introduce a new architecture for deep learning on point clouds, which is able to generate new points conditioned on their context. This architecture allows for large-scale point cloud data to be leveraged for representation learning. In the remainder of the paper, we describe these contributions in detail. We invite the community to use these benchmarks to test their model’s video understanding capabilities.

2. Related Work

Learning to persist objects through occlusions has been a long-standing challenge in computer vision [20, 36, 37]. In recent years, researchers have combined modern deep learned features with a variety of approaches to track through occlusions. These include classical Kalman filtering or linear extrapolation [25], 2D recurrent neural networks [57], and more explicit reasoning mechanisms [49]. Our approach tackles the problem in a more holistic manner, drawing on improvements in point cloud modeling, neural fields, and attention mechanisms. We briefly recap relevant work from each area.

Point cloud modeling. Earlier work on representing point clouds with deep networks is based on 2D projection [11, 24, 54] or 3D voxelization [33, 51]. These methods capitalize on the success of 2D and 3D convolutions in image and video understanding by preprocessing input point clouds into 2D or 3D grids. PointNet [44] proposed to use point-wise MLPs and pooling layers to compute permutation-invariant point cloud representations. PointNet was subsequently extended to allow for hierarchical features to better model local geometric structures [45], and combined with the idea of voxelization to create a highly efficient point cloud encoder [26]. More recently, researchers have begun to apply transformer attention mechanisms that were first found to be valuable in the language domain [58]

to encode point clouds [27, 28, 64]. To address the quadratic complexity in attention computation applied to large-scale point clouds, the Point Transformer [69] replaced global attention with local vector attention and introduced relative position encoding. We adopt the Point Transformer as our feature encoder backbone because of its efficiency and performance for various point cloud tasks.

Point cloud tasks. Our goal is somewhat similar to that of point completion networks [21, 55, 60, 62, 68], although these works typically operate on a per-object basis and address only *self*-occlusions or amodal completion. In contrast, we aim to reconstruct entire scenes and address the fundamental challenge of occlusions more generally. Because existing 4D architectures [13, 29, 56] lack a mechanism to efficiently create new points, they have not been demonstrated to be capable of dynamic scene completion. For example, in 4D panoptic LiDAR segmentation [3], the goal is to jointly tackle semantic and instance segmentation in 3D space over time. While our work addresses related tasks, we wish to be able to model not just the visible, but also the *occluded* parts of the scene, by drawing on past observations or priors. This is especially valuable when spatial inputs are sparse, as they often are in LiDAR applications.

Neural fields. Neural implicit functions have become very popular for 3D representation in recent years [12, 48, 50, 67], building on the seminal work of Neural Radiance Fields (NeRF) [34] and neural implicit surface modeling [39]. The basic idea of NeRF is to learn to represent a scene using a fully connected deep neural network, whose inputs are a 3D point and viewing direction and whose outputs are an estimated color and volume density. This is attractive because it avoids the need to discretize space and can encode a scene more efficiently and richly than traditional representations such as meshes or voxels, which themselves can be extracted from the implicit model. Numerous efforts have been made to extend NeRF to dynamic scenes [16, 17, 40, 43, 63], but in addition to requiring per-scene retraining, occlusions are typically explicitly ignored by applying losses over the non-occluded scene only.

Transformers in vision. The attention mechanism introduced in [4, 58] has been applied with great success to computer vision [10, 14, 53, 65]. Recently, architectures that are built solely with self-attention as computational units have started to perform on par with or better than convolutional networks as generic feature extractors [38] in standard vision tasks such as object detection and segmentation [30, 46, 47] and point cloud-based detection [35]. The role of cross-attention has also been extended as a mechanism for sensor fusion. DETR3D [61] extends DETR [8] by computing the keys and values from multi-view images. Recently, Perceiver [22] showed that asymmetric attention mechanisms can distill inputs from multiple modalities (*i.e.* vision or point clouds) into robust latent representations.

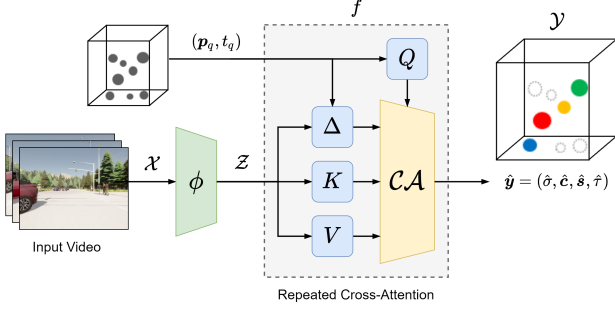


Figure 2. **Neural Architecture** – The encoder ϕ is a point transformer that featurizes the input point cloud video $\mathcal{X}$ using self-attention to produce $\mathcal{Z}$. The implicit representation f , conditioned on $\mathcal{Z}$, incorporates cross-attention blocks to contextualize the query points $(\mathbf{p}_q, t_q)$ and create the desired output features $\hat{\mathbf{y}}$ for that location and time within the scene.

3. Approach

We introduce the new task of 4D dynamic scene completion from posed monocular RGB-D video input. Let $\mathcal{X} = \{(\mathbf{p}_i, t_i, \mathbf{x}_i)\}$ be a point cloud video captured from a single camera view.² Each discrete point $(\mathbf{p}_i, t_i, \mathbf{x}_i)$ has a spatial position $\mathbf{p}_i \in \mathbb{R}^3$, a time $t_i \in \mathbb{R}$, and an RGB color $\mathbf{x}_i \in \mathbb{R}^3$ where the subscript i indicates the index. This information can be obtained realistically using a regular camera coupled with either a depth camera or a LiDAR sensor, aggregating data over multiple frames. Note that the input point cloud is only a partial scan, and consequently there are missing points due to occlusions or other noise, which makes this a challenging task. Our goal is to learn a mapping from $\mathcal{X}$ to a complete point cloud $\mathcal{Y} = \{(\mathbf{p}_j, t_j, \mathbf{y}_j)\}$ that densely encodes the full spacetime volume. The output vector $\mathbf{y}_j \in \mathbb{R}^d$ encodes any labels that we want to predict, such as color or semantic category.

3.1. Model

Point clouds are often treated as discrete, which causes them to have an irregular structure that makes traditional deep representation learning on them difficult. In order for our model to learn to persist points after they become occluded, we need a mechanism to create new points that have not been observed.

We will model the output point cloud as continuous, which allows us to compactly parameterize all the putative points across the 4D spacetime volume. Let $(\mathbf{p}_q, t_q) \in \mathbb{R}^4$ be a continuous spacetime query coordinate. Our model estimates the features $\hat{\mathbf{y}}$ located at $(\mathbf{p}_q, t_q)$, which may be occluded, with the decomposition:

$$\hat{\mathbf{y}}(\mathbf{p}_q, t_q) = f(\mathbf{p}_q, t_q; \phi(\mathcal{X})) \quad (1)$$

²A point cloud video assumes known camera parameters to deproject the RGB + depth information into some canonical coordinate system.

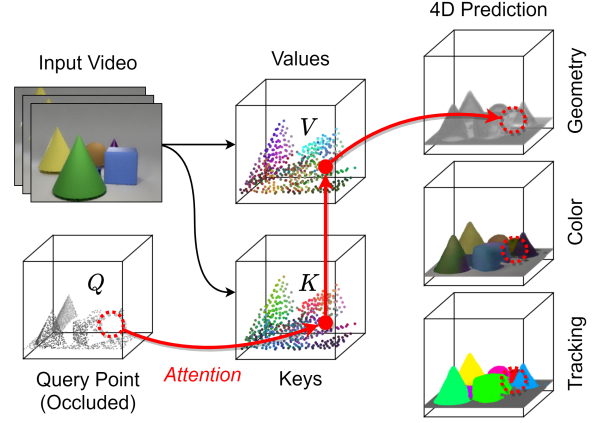


Figure 3. **Key Idea** – Given a query point in 4D, the model learns to attend to keys and values extracted from the input video. When the query point corresponds to world coordinates that are occluded, the attention mechanism will learn to find the object when it was not yet occluded. When the world coordinates correspond to empty space, the model instead learns to predict a low occupancy $\hat{y}_\sigma \approx 0$. The network is flexible, and we can train the same model to produce point clouds for many different video tasks requiring object permanence.

where ϕ is a feature extractor and f is our continuous representation. There are many possible choices for $\phi(\mathcal{X})$ [44, 45, 68], and we use the architecture from the Point Transformer network [69], which produces contextualized features for every point in the (subsamped) input.

The model is able to continuously predict a representation $\hat{\mathbf{y}}$ for the entire spacetime volume, shown in Figure 2. We can train $\hat{\mathbf{y}}$ for many different point cloud tasks, providing us the flexibility to predict, for example, geometry, semantic information, color, or object identity.

Our model uses a continuous representation similar to methods in neural rendering and computer graphics [16, 17, 40, 43, 63], which also enjoy significant computational advantages from the compact scene representation. However, our approach operates on point clouds instead of signed distance functions or radiance fields, which allows us to train and apply our model for many tasks besides view synthesis. Furthermore, our approach is conditioned on a set of frames in a dynamic point cloud video, enabling the model to learn a rich spatiotemporal representation for occluded objects.

3.2. Point Attention

Given the query coordinate $(\mathbf{p}_q, t_q)$, we need to estimate the contents at that spatiotemporal location. However, in a video with occlusions, the contextual evidence for those contents might be both spatially and temporally far away.

We introduce a cross-attention layer that uses the query coordinates to attend to the input video in order to generate this prediction. We illustrate this process in Figure 3. Typi-

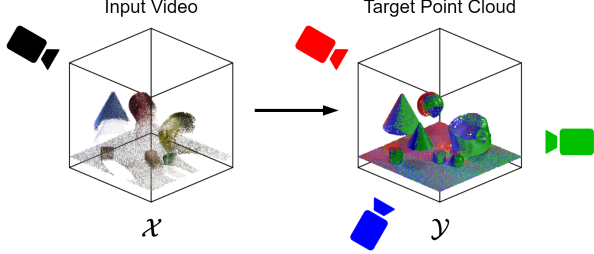


Figure 4. **Learning About Occlusions** – When an occlusion occurs in the input video, it is typically still visible from other view-points. For example, the smaller cone becomes fully occluded in the last input frame, but is revealed again by the ground truth. Our approach capitalizes on this natural clue to provide geometrically consistent multi-view self-supervision to the model, thus distilling the notion that objects persist across space and time.

cally, attention works by using a query to attend to keys and retrieve the relevant values. In our case, we will operate on the featurized point cloud $\mathcal{Z} = \phi(\mathcal{X}) = \{(\mathbf{p}_i, \boldsymbol{\alpha}_i)\}$. We form the keys $\mathbf{K}$ and values $\mathbf{V}$ from $\boldsymbol{\alpha}_i$, and the relative positional encodings Δ from $\mathbf{p}_q - \mathbf{p}_i$. We form the query $\mathbf{Q}$ from $\mathbf{p}_q$ and t_q .

Our layer can be recursively stacked, allowing the model to build increasingly rich representations of the scene. We implement the above vector cross-attention strategy through the following calculations, inspired by [69]:

$$\mathbf{Q}_q = \mathbf{w}_Q^\top \beta_{q,n} \quad (\text{queries}) \quad (2)$$

$$\mathbf{V}_i = \mathbf{w}_V^\top \boldsymbol{\alpha}_i \quad (\text{values}) \quad (3)$$

$$\mathbf{K}_i = \mathbf{w}_K^\top \boldsymbol{\alpha}_i \quad (\text{keys}) \quad (4)$$

$$\Delta_{q,i} = \mathbf{w}_\Delta^\top (\mathbf{p}_i - \mathbf{p}_q) \quad (\text{positions}) \quad (5)$$

where $\beta_{q,n}$ is a feature vector that encodes the features for query point $(\mathbf{p}_q, t_q)$. The base case is $\beta_{q,0} = \text{MLP}(\mathbf{p}_q, t_q)$, and as we stack the cross-attention block, it will be iteratively refined with:

$$\beta_{q,n+1} = \sum_{i \in \mathcal{M}(q)} \rho(\gamma(\mathbf{Q}_q - \mathbf{K}_i + \Delta_{q,i})) \odot (\mathbf{V}_i + \Delta_{q,i}) \quad (6)$$

where $\mathcal{M}(q)$ is a set of nearest neighbors within $\phi(\mathcal{X})$ around $\mathbf{p}_q$, ρ is the softmax operation for normalization, γ is a mapping MLP that produces the attention weights, and $\odot$ is an element-wise product to represent per-channel feature modulation [69].

We apply the operation in Equation (6) twice, meaning we terminate the recursion at $\beta_{q,2}$. This produces a feature vector that describes the contents at the query location, which we decode into the predicted labels. Finally, we use an MLP to map $\beta_{q,2}$ to $\hat{\mathbf{y}}(\mathbf{p}_q, t_q)$.

3.3. Learning and Supervision

We train the model for 4D dynamic scene completion. Given several camera views of a scene, we assume known camera parameters to deproject their recordings into point clouds. We select one camera view to be the input view, which creates $\mathcal{X}$. To form the target $\mathcal{Y}$, we use the point cloud that merges all the camera views together. We train the model to predict the multi-view point cloud $\mathcal{Y}$ from the single-view point cloud $\mathcal{X}$, illustrated in Figure 4.

Due to the efficiency of our representation, we can train the model end-to-end for large spacetime volumes on standard GPU hardware. We minimize the loss function:

$$\min_{f, \phi} \mathbb{E}_{(\mathcal{X}, \mathcal{Y})} \left[\sum_{(\mathbf{p}_q, t_q, \mathbf{y}_q) \in \mathcal{Y} \cup \mathcal{N}} \mathcal{L}(\hat{\mathbf{y}}(\mathbf{p}_q, t_q), \mathbf{y}_q) \right] \quad (7)$$

where $\mathcal{N}$ is a set of negative points randomly sampled uniformly from $\mathbb{R}^4$. Since the training data $\mathcal{Y}$ only contains solid points, the negative points cause the model to learn to distinguish which regions are empty space.

3.4. Tasks

Our framework is able to learn to reveal occlusions for several different tasks on point clouds. For every query point, the model produces a vector $\hat{\mathbf{y}}_i \in \mathbb{R}^d$, and we can supervise different dimensions of $\hat{\mathbf{y}}_i$ for various tasks. We select the loss function $\mathcal{L}$ depending on the dataset and task. We describe several options for the loss terms below.

Geometry completion distinguishes solid objects ($\sigma = 1$) from free space ($\sigma = 0$) within the scene, where the ground truth occupancy σ is inferred for every query point by thresholding its proximity to the target point cloud. Denoting $\hat{y}_\sigma$ as the relevant dimension of the $\hat{\mathbf{y}}$ vector, we apply a standard binary cross-entropy comparison as follows: $\mathcal{L}_\sigma = \mathcal{L}_{BCE}(\hat{y}_\sigma, \sigma)$.

Visual reconstruction means that, in addition to completing the missing regions, the model must also predict a color $\hat{\mathbf{y}}_c$ in RGB space. For the loss function, we use the $\mathcal{L}_1$ -distance between the relevant output dimensions and the target $\mathbf{c}$: $\mathcal{L}_c = \|\hat{\mathbf{y}}_c - \mathbf{c}\|_1$.

Semantic segmentation classifies every query point into S possible categories. The output is supervised with a cross-entropy loss between the predicted categories $\hat{\mathbf{y}}_s$ and the ground truth semantic label s : $\mathcal{L}_s = \mathcal{L}_{CE}(\hat{\mathbf{y}}_s, s)$.

Instance tracking tasks the model with localizing an object, even through total occlusions, that was highlighted with a mask in only the first frame.³ To do this, we add an extra dimension τ_i to the input point cloud $\mathcal{X}$, that indicates which points belong to the object of interest. We then train

³This is similar to most *semi-supervised video object segmentation* setups [6, 59], but in 3D space instead. Note that the object may be partially, but not completely, occluded at the beginning of the video for this to work.

the model to propagate this indicator throughout the rest of the video, where $\hat{y}_\tau$ is the relevant dimension in the output. We use the binary cross-entropy loss between the tracking flag $\hat{y}_\tau$ and τ : $\mathcal{L}_\tau = \mathcal{L}_{BCE}(\hat{y}_\tau, \tau)$.

These four loss terms can be linearly combined to form the overall objective:

$$\mathcal{L} = \lambda_\sigma \mathcal{L}_\sigma + \lambda_c \mathcal{L}_c + \lambda_s \mathcal{L}_s + \lambda_\tau \mathcal{L}_\tau \quad (8)$$

Geometry completion $\mathcal{L}_\sigma$ is supervised in all of space-time, while the other three loss terms $\mathcal{L}_c$, $\mathcal{L}_s$, $\mathcal{L}_\tau$ are applied in solid regions only.

3.5. Inference

After learning, we will be able to estimate a continuous representation of a point cloud from a video. For many applications, we need a sampling procedure to convert the continuous cloud into a discrete cloud. Depending on our choice of sampling technique, we can construct arbitrarily detailed point clouds at test time.

Since the target is unknown at test time, we sample query coordinates (p_q, t_q) uniformly at random within a 4D spacetime volume of interest. We generate discrete point clouds by filtering predictions according to solidity, only retaining a query point whenever the predicted occupancy is above some threshold, *i.e.* $\hat{y}_\sigma \geq \sigma_T$.

For visualization purposes, we can also convert the predictions to scene meshes. The surface $\mathcal{S}$ of a mesh at time t is implicitly defined as the zero-level set of the predicted occupancy $\hat{\sigma}$ relative to the threshold σ_T , *i.e.* $\mathcal{S} = \{x \in \mathbb{R}^3 \mid \hat{y}_\sigma(x, t) = \sigma_T\}$, where $\sigma_T = 0.5$.

After sampling a point cloud, or a mesh via the marching cubes algorithm, we colorize it by retrieving either the predicted color $\hat{y}_c$, the semantic category $\hat{y}_s$, or the tracking flag $\hat{y}_\tau$, associated with every coordinate.

3.6. Implementation Details

The feature encoder ϕ interleaves 4 self-attention layers with 3 down transition modules [69] to generate the featurized point cloud $\mathcal{Z}$ from $\mathcal{X}$. The continuous representation f , conditioned on $\mathcal{Z}$, accepts arbitrary 4D query coordinates (p_q, t_q) as input, applies Fourier encoding [34], and interleaves 6 residual MLP blocks [67] with 2 cross-attention layers to produce $\hat{y}$.

We feed in $T = 12$ frames with $|\mathcal{X}| = 14,336$ points in total, and train the model to predict the last $U = 4$ frames, such that the first $T - U = 8$ frames serve as an opportunity to aggregate and process spatiotemporal context. More details can be found in the supplementary material.

4. Datasets

In order to train and evaluate our model, particularly in terms of its ability to handle occlusions, we require

multi-view RGB-D video from highly cluttered scenes. To this end, we contribute two high-quality synthetic datasets, shown in Figs. 5 and 6. Brief descriptions are provided below, with further details in the supplementary material.

4.1. GREATER

We extend CATER [18] (which is in turn based on CLEVR [23]) mainly in order to increase the degree of occlusions, and call our proposed dataset GREATER. Each scene in GREATER contains 8 to 12 cubes, cones, cylinders and spheres that move around, occluding one another in random ways. Partial and complete occlusions are happening constantly to the input view, which are only revealed by the target point clouds, allowing for effective learning and benchmarking of our model. We capture 7,000 scenes lasting 12 seconds each, with data captured from 3 random views spaced at least 45° apart horizontally, and a train/val/test split of 80%/10%/10%. On the GREATER dataset, we train our model to predict geometry, color, and tracking.

4.2. CARLA

While GREATER already exhibits many non-trivial scene configurations and movement patterns, it may be desirable to apply 4D video completion within more realistic environments as well. Since object permanence is paramount for situational awareness in the context of driving and traffic scenarios, we employ the state-of-the-art driving simulator CARLA [15] to generate a dataset of complex, dynamic road scenes. We sample 500 scenes lasting 100 seconds each, with data captured from 4 fixed views, and a train/val/test split of 80%/8%/12%. The scenes cover a wide variety of different towns, vehicles, pedestrians, traffic scenarios, and weather conditions. On the CARLA dataset, we teach our model to perform geometric as well as semantic scene completion.

5. Experiments

To display the generality of our framework, we test it on a variety of tasks across different datasets. Crucially, all the tasks can be trained end-to-end simultaneously in the same model. We set $(\lambda_\sigma, \lambda_c, \lambda_s, \lambda_\tau) = (1, 1, 0, 1)$ for GREATER and $(\lambda_\sigma, \lambda_c, \lambda_s, \lambda_\tau) = (1, 0, 0.6, 0)$ for CARLA.

5.1. Evaluation Metrics

We evaluate models using the Chamfer Distance (CD) metric between the predicted point cloud $\hat{\mathcal{Y}}$ and target point cloud $\mathcal{Y}$:

$$CD(\hat{\mathcal{Y}}, \mathcal{Y}) = \frac{1}{|\hat{\mathcal{Y}}|} \sum_{i \in \hat{\mathcal{Y}}} \min_{j \in \mathcal{Y}} \|p_i - p_j\|_2 + \frac{1}{|\mathcal{Y}|} \sum_{j \in \mathcal{Y}} \min_{i \in \hat{\mathcal{Y}}} \|p_j - p_i\|_2 \quad (9)$$

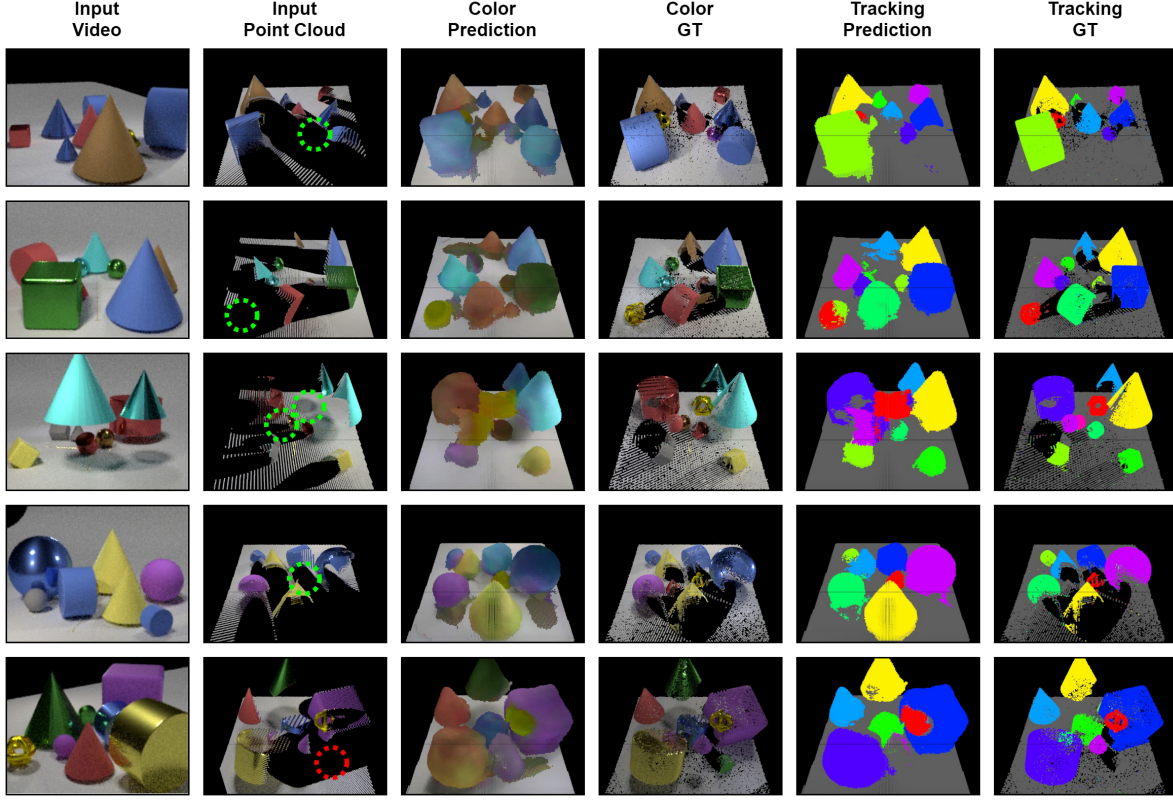


Figure 5. **Results for GREATER** – We show inputs, predictions, and ground truths. Our model receives color point clouds as input (second column), and we show the corresponding video frame in column one as reference. The third column represents *both* the geometry reconstruction and color prediction tasks. We note how the model is able to (1) perform scene completion by filling in partially observed objects, *i.e.* resolve amodal completion, and even (2) recover totally occluded objects, including when there are multiple occurring at once. For *total* occlusions, we circle the corresponding locations in the input for **true positives** in green and **false negatives** in red. While we show only the last frames in this figure, the model predicts the scene at different time steps, capturing scene dynamics.

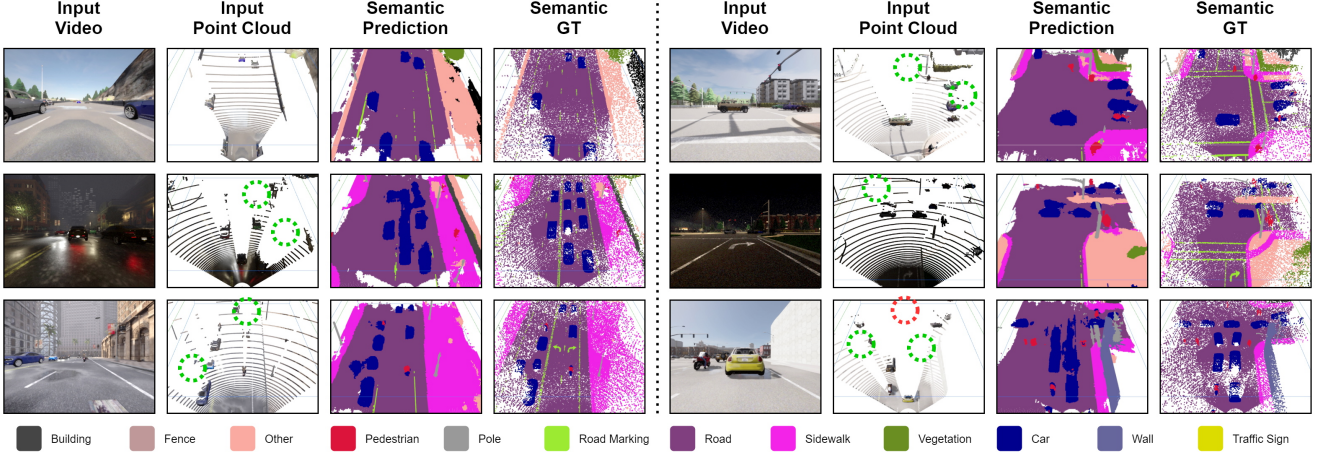


Figure 6. **Results for CARLA** – We show inputs, predictions and ground truths. Our model receives the point cloud video whose last frame is depicted in the second column, and predicts scene occupancy and semantic completion data for every sampled query point. Considering the limited input information, our model is capable of reconstructing the whole scene with high accuracy. Just as in Figure 5, all inputs and outputs are 4D meaning that they actually consist of multiple frames – please see our [webpage](#) for animated visualizations.

Method	Geometry		Tracking	
	All	Occ.	Avg. Inst.	Occ.
No local features	0.78	0.73	4.50	4.86
No time	0.26	0.49	1.59	4.25
No self-attention	0.26	0.37	1.12	1.58
No cross-attention	0.32	0.41	1.21	1.73
No attention	0.40	0.48	1.42	2.00
Copy input view	0.48	1.92	-	-
Assume stationary	-	-	2.34	3.43
PCN [68]	0.59	0.97	-	-
Ours	0.22	0.33	1.05	1.32

Table 1. **Results for GREATER – geometry completion and instance tracking tasks.** We report the Chamfer Distance (lower is better). In addition to outperforming all ablations and baselines on both tasks, our model predicts occluded objects nearly as well as visible objects.

For geometry completion, we initially consider all points, but wish to specifically study occlusions as well. To that end, we filter all points by whether they belong to an occluded instance or not, which we can approximate by comparing different views with each other. If the filtered output point cloud is empty (which typically corresponds to false negatives), we substitute the prediction for a single point at the center of the scene, as the CD would become undefined otherwise.

For instance tracking in GREATER, we track one object at a time, and merge the resulting predictions at test time. Concretely, we obtain multiple tracks by assigning the instance tag with the most confident score $\hat{y}_\tau$ to each point, but only if $\hat{y}_\tau \geq 0.5$. Then, for every instance tag, we calculate the CD between only its corresponding predicted points and the ground truth object points, and subsequently average this value over all instances within a scene. We also report the average over occluded objects only.

For semantic segmentation in CARLA, we use a similar workflow as for tracking, but average over all categories instead of instances. We study two important classes (pedestrians and vehicles) separately, which implies filtering both the predictions and targets by whether their ground truth semantic categories belong to those respective classes before reporting the CD values. Additionally, we filter for *occluded* pedestrians and vehicles. In both cases, we average over all instances per scene such that every pedestrian or car is treated equally.

5.2. Ablations and Baselines

Ablations. To show how various architectural choices affect our model’s performance, we perform ablations to its four main components by: (1) removing local features from f ; (2) removing the temporal dimension; (3) removing self-attention from ϕ ; (4) removing cross-attention from f ; (5) combining (3) and (4). Ablation (1) implies that instead of

Method	Geometry All	Semantic Segmentation				
		Avg. Cls.	Ped.	Veh.	Occ. Ped.	Occ. Veh.
No local features	1.12	15.23	19.26	6.81	19.95	6.97
No time	0.55	6.76	10.16	5.51	15.56	10.99
No self-attention	0.50	6.60	5.19	3.29	5.98	4.73
No cross-attention	0.73	7.11	7.57	4.22	11.33	7.55
No attention	0.71	9.21	9.68	4.61	13.66	7.51
Copy input view	1.39	-	-	-	-	-
PCN [68]	11.79	-	-	-	-	-
Ours	0.47	5.82	4.76	3.42	5.71	4.14

Table 2. **Results for CARLA – geometry completion and semantic completion tasks.** We report the Chamfer Distance (lower is better). Our model significantly outperforms almost all baselines and ablations, especially for occluded pedestrians (“Ped.”) and vehicles (“Veh.”).

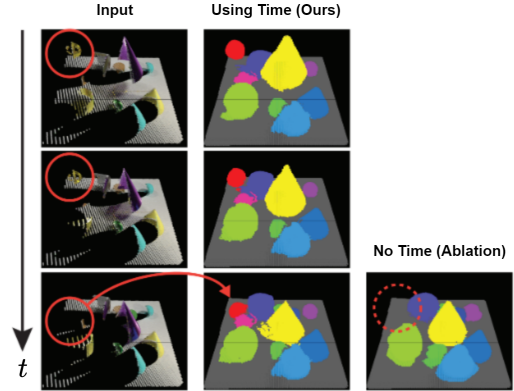


Figure 7. **Importance of Time** – In this figure, we show a video where an occlusion takes place (see yellow orb, highlighted by the red circle), and tracking predictions of our model. Because our model takes video as input, it is capable of tracking the orb through the occlusion. The single-frame baseline (“no time”) cannot recover the occluded orb, as it does not have access to temporal context. This example shows that (1) temporal context is important, and (2) our model learns to use it.

conditioning on $\mathcal{Z}$, we only pass a global embedding (that is the average of the features over all points in $\mathcal{Z}$) from ϕ to f . In (2), the model is no longer burdened with predicting a 4D dynamic representation consisting of multiple frames, and the task becomes 3D scene completion instead – given a single frame, predict a single frame. For (3), we replace self-attention layers in the point transformer with a simple point-wise linear projection. For (4), since f cannot attend to $\mathcal{Z}$ anymore, we feed in the nearest neighbor in $\mathcal{Z}$ of every query point to f along with the query point itself.

Baselines. For our primary scene geometry reconstruction task, we adapt Point Completion Network (PCN) [68] to our setup. Additionally, we evaluate a ‘Copy input view’ baseline, where the prediction is simply the input point cloud

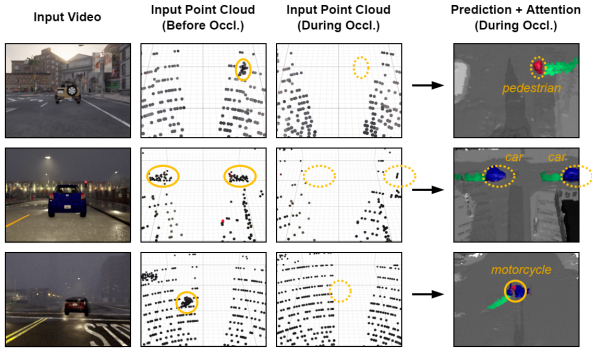


Figure 8. **Visualizing Attention** – Why did the model predict the **occluded object(s)**? By backtracking neuron activations through all cross-attention and self-attention layers, we see that a mechanism of temporal correspondence emerges. In this example, the **attention weights** highlight input points that represent the trajectory of the object(s) over time, which suggests that our model implicitly learns to track them in order to succeed at 4D scene completion.

that the model sees, *i.e.* the identity operation. Comparison with this baseline shows the benefits of our approach in revealing occlusions. Finally, for tracking, we evaluate the baseline where the marked instance is propagated but remains stationary after the first frame, which is also the only time that the model sees its mask.

5.3. Quantitative Results

See Tables 1 and 2. Occlusion metrics (“Occ.”) are for objects that are more than 80% occluded, as inferred from counting the number of points per instance that are visible from each view. Our non-ablated model consistently outperforms most baselines and ablations with significant margins. In particular, these results demonstrate that incorporating attention mechanisms is clearly beneficial for handling occlusions and performing spatiotemporal inpainting, suggesting that a robust notion of object permanence was successfully learned.

Although the ablation without time succeeds at reconstructing a 3D snapshot of the scene with relatively good quality, it is significantly worse at predicting occluded objects such as vehicles or pedestrians in CARLA, which is a critical aspect of interpreting traffic scenes. Figure 7 further demonstrates that temporal context is essential to understand scene dynamics. Moreover, providing contextualized local features is also vital for the performance of the model.

5.4. Visualizations

In this section, we visualize the inputs and predictions of our model, for both datasets. Then, we focus our attention on how our model deals with two key challenges the task presents: occlusions and uncertainty.

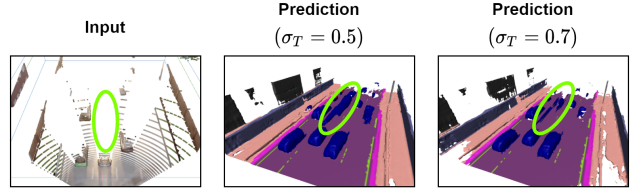


Figure 9. **Interpreting Uncertainty** – Because it is often hard to perceive what is ahead of the car in front of you, some regions of the input can remain unseen in the LiDAR point cloud (shown on the left) throughout the entire input video, causing artefacts such as a ‘long car’ to crop up. By varying the occupancy threshold σ_T of the implicit surface, we can control the degree of certainty at which to visualize the prediction. For example, in the above example the model is found to be less certain about the presence of the ‘long car’ relative to the rest of the scene (which remains reconstructed accurately), as indicated by the artefact’s disappearance with increasing σ_T .

Figures 5 and 6 show our model predictions for the GREATER and CARLA datasets. In both cases, the model is capable of simultaneously performing geometry completion along with other prediction tasks such as visual reconstruction, instance tracking, or semantic segmentation. Note that geometry completion is a prerequisite to solving any other prediction task, since other tasks also require knowledge of objects that are *not visible* in the input frame.

Our model is capable of completing the scene with great detail, even when presented with a limited density of input points. Specifically, when trained on CARLA the model is capable of reconstructing—and predicting class information about—relatively small objects such as street poles (gray), pedestrians (red), or traffic signs (yellow). It does so for different temporal steps, and even when there are occlusions. We observe that the model trained on CARLA does sometimes struggle in hard cases that involve objects moving across long-term occlusions (which presents a limitation and opportunity for future work), but it usually generates an accurate, complete reconstruction in most other scenarios, especially when just amodal completion is involved.

In Figure 7 we visualize how our 4D model is capable of exploiting temporal context to track through occlusions, unlike the “no time” baseline. We also show that the model is 4D in the sense that it can represent points for different time steps.

In Figure 8, we visualize the attention of the model during occlusions in order to understand the mechanism it uses to resolve them. Specifically, we adapt attention rollout [1] to our architecture, and visualize the input points (across time) that contribute the most to the specific output class we want to analyze. We show results for a pedestrian, two cars, and a motorcycle. In all cases, the attention for the studied object is focused on its past trajectory, meaning that

the model implicitly tracks objects through time in order to make predictions about their future.

Finally, the reader may have noticed some “long cars” in Figure 6. These occur in places that remain occluded throughout the video, and Figure 9 illustrates that our model appears to encode uncertainty with respect to the scene contents to some extent. This allows us to control the degree of certainty we want to visualize the prediction at, and this parameter can optionally be tuned depending on the downstream application.

Please refer to the supplementary material for extra, non-cherry-picked visualizations that include more baselines.

6. Discussion

We introduce the task of 4D dynamic scene completion, along with two datasets for understanding occlusions, and showcase a continuous representation that incorporates cross-attention as an initial attempt toward solving this challenge. We believe these techniques and benchmarks will be useful in the context of scene completion, spatiotemporal inpainting, and object permanence.

Acknowledgements: This research is based on work supported by Toyota Research Institute, the NSF CAREER Award #2046910, and the DARPA MCS program under Federal Agreement No. N660011924032. DS is supported by the Microsoft PhD Fellowship. The views and conclusions contained herein are those of the authors and should not be interpreted as necessarily representing the official policies, either expressed or implied, of the sponsors.

References

- [1] Samira Abnar and Willem Zuidema. Quantifying attention flow in transformers. In *ACL*, 2020. 8
- [2] Andréa Aguiar and Renée Baillargeon. 2.5-month-old infants’ reasoning about when objects should and should not be occluded. *Cognitive psychology*, 39(2):116–157, 1999. 1
- [3] Mehmet Aygun, Aljosa Osep, Mark Weber, Maxim Maximov, Cyrill Stachniss, Jens Behley, and Laura Leal-Taixé. 4d panoptic lidar segmentation. In *CVPR*, 2021. 2
- [4] Dzmitry Bahdanau, Kyunghyun Cho, and Yoshua Bengio. Neural machine translation by jointly learning to align and translate. *arXiv preprint arXiv:1409.0473*, 2014. 2
- [5] Renee Baillargeon, Elizabeth S Spelke, and Stanley Wasserman. Object permanence in five-month-old infants. *Cognition*, 20(3):191–208, 1985. 1
- [6] Sergi Caelles, Alberto Montes, Kevis-Kokitsi Maninis, Yuhua Chen, Luc Van Gool, Federico Perazzi, and Jordi Pont-Tuset. The 2018 davis challenge on video object segmentation. *arXiv preprint arXiv:1803.00557*, 2018. 4
- [7] Holger Caesar, Varun Bankiti, Alex H Lang, Sourabh Vora, Venice Erin Liong, Qiang Xu, Anush Krishnan, Yu Pan, Giancarlo Baldan, and Oscar Beijbom. nuScenes: A multi-modal dataset for autonomous driving. In *CVPR*, 2020. 1
- [8] Nicolas Carion, Francisco Massa, Gabriel Synnaeve, Nicolas Usunier, Alexander Kirillov, and Sergey Zagoruyko. End-to-end object detection with transformers. In *ECCV*, 2020. 2
- [9] Ming-Fang Chang, John Lambert, Patsorn Sangkloy, Jagjeet Singh, Slawomir Bak, Andrew Hartnett, De Wang, Peter Carr, Simon Lucey, Deva Ramanan, et al. Argoverse: 3d tracking and forecasting with rich maps. In *CVPR*, 2019. 1
- [10] Xiangning Chen, Cho-Jui Hsieh, and Boqing Gong. When vision transformers outperform resnets without pre-training or strong data augmentations. *arXiv preprint arXiv:2106.01548*, 2021. 2
- [11] Xiaozhi Chen, Huimin Ma, Ji Wan, Bo Li, and Tian Xia. Multi-view 3d object detection network for autonomous driving. In *CVPR*, 2017. 2
- [12] Julian Chibane, Aymen Mir, and Gerard Pons-Moll. Neural unsigned distance fields for implicit function learning. In *NeurIPS*, 2020. 2
- [13] Christopher Choy, JunYoung Gwak, and Silvio Savarese. 4d spatio-temporal convnets: Minkowski convolutional neural networks. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 3075–3084, 2019. 2
- [14] Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, Jakob Uszkoreit, and Neil Houlsby. An image is worth 16x16 words: Transformers for image recognition at scale. *ICLR*, 2021. 2
- [15] Alexey Dosovitskiy, German Ros, Felipe Codevilla, Antonio Lopez, and Vladlen Koltun. Carla: An open urban driving simulator. In *CoRL*. PMLR, 2017. 5
- [16] Yilun Du, Yinan Zhang, Hong-Xing Yu, Joshua B. Tenenbaum, and Jiajun Wu. Neural radiance flow for 4d view synthesis and video processing. In *ICCV*, 2021. 2, 3
- [17] Chen Gao, Ayush Saraf, Johannes Kopf, and Jia-Bin Huang. Dynamic view synthesis from dynamic monocular video. In *ICCV*, 2021. 2, 3
- [18] Rohit Girdhar and Deva Ramanan. CATER: A diagnostic dataset for Compositional Actions and Temporal Reasoning. In *ICLR*, 2020. 5, 3
- [19] Meng-Hao Guo, Jun-Xiong Cai, Zheng-Ning Liu, Tai-Jiang Mu, Ralph R Martin, and Shi-Min Hu. PCT: Point cloud transformer. *Computational Visual Media*, 7(2):187–199, 2021. 1
- [20] Yan Huang and Irfan Essa. Tracking multiple objects through occlusions. In *CVPR*, 2005. 2
- [21] Zitian Huang, Yikuan Yu, Jiawen Xu, Feng Ni, and Xinyi Le. Pf-net: Point fractal network for 3d point cloud completion. In *CVPR*, 2020. 2
- [22] Andrew Jaegle, Felix Gimeno, Andrew Brock, Andrew Zisserman, Oriol Vinyals, and Joao Carreira. Perceiver: General perception with iterative attention. In *ICML*, 2021. 2
- [23] Justin Johnson, Bharath Hariharan, Laurens Van Der Maaten, Li Fei-Fei, C Lawrence Zitnick, and Ross Girshick. CLEVR: A diagnostic dataset for compositional language and elementary visual reasoning. In *CVPR*, 2017. 5

- [24] Asako Kanezaki, Yasuyuki Matsushita, and Yoshifumi Nishida. Rotationnet for joint object categorization and unsupervised pose estimation from multi-view images. *TPAMI*, 43(1):269–283, 2021. 2
- [25] Tarasha Khurana, Achal Dave, and Deva Ramanan. Detecting invisible people. In *ICCV*, 2021. 2
- [26] Alex H. Lang, Sourabh Vora, Holger Caesar, Lubing Zhou, Jiong Yang, and Oscar Beijbom. Pointpillars: Fast encoders for object detection from point clouds. In *CVPR*, 2019. 2
- [27] Juho Lee, Yoonho Lee, Jungtaek Kim, Adam Kosiorek, Seungjin Choi, and Yee Whye Teh. Set transformer: A framework for attention-based permutation-invariant neural networks. In *ICML*, 2019. 2
- [28] Xin Hai Liu, Zhizhong Han, Yu-Shen Liu, and Matthias Zwicker. Point2sequence: Learning the shape representation of 3d point clouds with an attention-based sequence to sequence network. In *AAAI*, 2019. 2
- [29] Xingyu Liu, Mengyuan Yan, and Jeannette Bohg. Meteor-net: Deep learning on dynamic 3d point cloud sequences. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 9246–9255, 2019. 2
- [30] Ze Liu, Yutong Lin, Yue Cao, Han Hu, Yixuan Wei, Zheng Zhang, Stephen Lin, and Baining Guo. Swin transformer: Hierarchical vision transformer using shifted windows. In *ICCV*, 2021. 2
- [31] Ilya Loshchilov and Frank Hutter. Decoupled weight decay regularization. In *ICLR*, 2019. 2
- [32] Lori Marino. Thinking chickens: a review of cognition, emotion, and behavior in the domestic chicken. *Animal Cognition*, 20(2):127–147, 2017. 1
- [33] Daniel Maturana and Sebastian Scherer. Voxnet: A 3d convolutional neural network for real-time object recognition. In *IROS*, 2015. 2
- [34] 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. In *ECCV*, 2020. 2, 5, 3
- [35] Ishan Misra, Rohit Girdhar, and Armand Joulin. An End-to-End Transformer Model for 3D Object Detection. In *ICCV*, 2021. 2
- [36] Hieu Tat Nguyen, Marcel Worring, and Rein Van Den Boomgaard. Occlusion robust adaptive template tracking. In *ICCV*, 2001. 2
- [37] Jiyan Pan and Bo Hu. Robust occlusion handling in object tracking. In *CVPR*, 2007. 2
- [38] Zizheng Pan, Bohan Zhuang, Jing Liu, Haoyu He, and Jianfei Cai. Scalable vision transformers with hierarchical pooling. In *ICCV*, 2021. 2
- [39] Jeong Joon Park, Peter Florence, Julian Straub, Richard Newcombe, and Steven Lovegrove. Deepsdf: Learning continuous signed distance functions for shape representation. In *CVPR*, 2019. 2
- [40] Keunhong Park, Utkarsh Sinha, Jonathan T. Barron, Sofien Bouaziz, Dan B Goldman, Steven M. Seitz, and Ricardo Martin-Brualla. NeRFies: Deformable neural radiance fields. *ICCV*, 2021. 2, 3
- [41] Irene M Pepperberg and Mildred S Funk. Object permanence in four species of psittacine birds: An african grey parrot (*psittacus erithacus*), an illiger mini macaw (*ara maracana*), a parakeet (*melopsittacus undulatus*), and a cockatiel (*nymphicus hollandicus*). *Animal Learning & Behavior*, 18(1):97–108, 1990. 1
- [42] Jean Piaget. *The construction of reality in the child*, volume 82. Routledge, 2013. 1
- [43] Albert Pumarola, Enric Corona, Gerard Pons-Moll, and Francesc Moreno-Noguer. D-NeRF: Neural radiance fields for dynamic scenes. In *CVPR*, 2021. 2, 3
- [44] Charles R Qi, Hao Su, Kaichun Mo, and Leonidas J Guibas. Pointnet: Deep learning on point sets for 3d classification and segmentation. In *CVPR*, 2017. 1, 2, 3
- [45] Charles R Qi, Li Yi, Hao Su, and Leonidas J Guibas. Pointnet++: Deep hierarchical feature learning on point sets in a metric space. In *NeurIPS*, 2017. 2, 3
- [46] René Ranftl, Alexey Bochkovskiy, and Vladlen Koltun. Vision transformers for dense prediction. In *ICCV*, 2021. 2
- [47] René Ranftl, Katrin Lasinger, David Hafner, Konrad Schindler, and Vladlen Koltun. Towards robust monocular depth estimation: Mixing datasets for zero-shot cross-dataset transfer. *TPAMI*, 2020. 2
- [48] Christoph Rist, David Emmerichs, Markus Enzweiler, and Darius Gavrila. Semantic scene completion using local deep implicit functions on lidar data. *TPAMI*, 2021. 2
- [49] Aviv Shamsian, Ofri Kleinfeld, Amir Globerson, and Gal Chechik. Learning object permanence from video. In *ECCV*, 2020. 2
- [50] Vincent Sitzmann, Julien Martel, Alexander Bergman, David Lindell, and Gordon Wetzstein. Implicit neural representations with periodic activation functions. *NeurIPS*, 2020. 2
- [51] Shuran Song, Fisher Yu, Andy Zeng, Angel X Chang, Manolis Savva, and Thomas Funkhouser. Semantic scene completion from a single depth image. In *CVPR*, 2017. 2
- [52] Elizabeth S Spelke and Gretchen Van de Walle. Perceiving and reasoning about objects: Insights from infants. *Spatial representation: Problems in philosophy and psychology*, pages 132–161, 1993. 1
- [53] Andreas Steiner, Alexander Kolesnikov, Xiaohua Zhai, Ross Wightman, Jakob Uszkoreit, and Lucas Beyer. How to train your ViT? Data, augmentation, and regularization in vision transformers. *arXiv preprint arXiv:2106.10270*, 2021. 2
- [54] Hang Su, Subhransu Maji, Evangelos Kalogerakis, and Erik G. Learned-Miller. Multi-view convolutional neural networks for 3d shape recognition. In *ICCV*, 2015. 2
- [55] Lyne P Tchapmi, Vineet Kosaraju, Hamid Reza Tofighi, Ian Reid, and Silvio Savarese. Topnet: Structural point cloud decoder. In *CVPR*, 2019. 2
- [56] Hugues Thomas, Charles R Qi, Jean-Emmanuel Deschaud, Beatriz Marcotegui, François Goulette, and Leonidas J Guibas. Kpconv: Flexible and deformable convolution for point clouds. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 6411–6420, 2019. 2
- [57] Pavel Tokmakov, Jie Li, Wolfram Burgard, and Adrien Gaidon. Learning to track with object permanence. In *ICCV*, 2021. 2

- [58] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. In *NeurIPS*, 2017. 2
- [59] Qiang Wang, Li Zhang, Luca Bertinetto, Weiming Hu, and Philip HS Torr. Fast online object tracking and segmentation: A unifying approach. In *CVPR*, 2019. 4
- [60] Xiaogang Wang, Marcelo H Ang Jr, and Gim Hee Lee. Cascaded refinement network for point cloud completion. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 790–799, 2020. 2
- [61] Yue Wang, Vitor Guizilini, Tianyuan Zhang, Yilun Wang, Hang Zhao, , and Justin M. Solomon. Detr3d: 3d object detection from multi-view images via 3d-to-2d queries. In *CoRL*, 2021. 2
- [62] Xin Wen, Tianyang Li, Zhizhong Han, and Yu-Shen Liu. Point cloud completion by skip-attention network with hierarchical folding. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 1939–1948, 2020. 2
- [63] Wenqi Xian, Jia-Bin Huang, Johannes Kopf, and Changil Kim. Space-time neural irradiance fields for free-viewpoint video. In *CVPR*, 2021. 2, 3
- [64] Saining Xie, Sainan Liu, Zeyu Chen, and Zhuowen Tu. Attentional shapecontextnet for point cloud recognition. In *CVPR*, 2018. 2
- [65] Zhenjia Xu, Zhanpeng He, Jiajun Wu, and Shuran Song. Learning 3d dynamic scene representations for robot manipulation. In *CoRL*, 2020. 2
- [66] Guandao Yang, Xun Huang, Zekun Hao, Ming-Yu Liu, Serge Belongie, and Bharath Hariharan. Pointflow: 3d point cloud generation with continuous normalizing flows. In *ICCV*, 2019. 1
- [67] Alex Yu, Vickie Ye, Matthew Tancik, and Angjoo Kanazawa. PixelNeRF: Neural radiance fields from one or few images. In *CVPR*, 2021. 2, 5, 3
- [68] Wentao Yuan, Tejas Khot, David Held, Christoph Mertz, and Martial Hebert. PCN: Point completion network. In *3DV*, 2018. 1, 2, 3, 7, 4
- [69] Hengshuang Zhao, Li Jiang, Jiaya Jia, Philip HS Torr, and Vladlen Koltun. Point transformer. In *ICCV*, 2021. 1, 2, 3, 4, 5

Revealing Occlusions with 4D Neural Fields

Supplementary Material

A. Implementation Details

A.1. Architecture and Hyperparameters

We illustrate the full model architecture in Figure 10. The input point cloud video $\mathcal{X}$ is subsampled using a mixture of random point sampling and farthest point sampling to contain precisely 14,336 points with 7 to 8 features each. After the first MLP, the embedding size per point is 36, which becomes double at every down transition step, such that $\mathcal{Z}$ has 288 features per embedding. The encoder ϕ also averages all features into a 128-dimensional global embedding, which is also passed to f alongside the featurized point cloud $\mathcal{Z}$. The latent size of embeddings within f is 416 across all its components.

The self-attention block causes every point to attend to its 16 nearest neighbors [69]. The down transition block selects one third of all incoming points by means of farthest point sampling, and subsequently performs channel-wise max-pooling from its 12 nearest neighbors. The residual MLP block in f initially distills a weighted linear interpolation (based on Euclidean distance) of the 8 nearest neighbors in $\mathcal{Z}$ around the query point as a starting point, after which cross-attention has the chance to apply a *learned* interpolation of embeddings instead. The cross-attention block causes every query point to attend to its 14 nearest neighbors in the featurized input points $\mathcal{Z}$, *i.e.* $|\mathcal{M}(i)| = 14, \forall i$.

Using the AdamW optimizer [31], we train two separate models (one per dataset) over 20 epochs for GREATER, and 40 epochs for CARLA. Our model takes between 18 and 55 hours to train on two RTX A6000 GPUs, and dense inference across the entire spacetime cube takes roughly one minute. The initial learning rate is 0.001, but this drops by a factor 2.5 at progress rates of 40%, 60%, and 80%.

A.2. Point Sampling

During training, within every frame we sample 7,168 solid query points ($\sigma = 1$) and 10,752 free space (air) query points ($\sigma = 0$). The air points are uniformly randomly sampled within the output box of interest, except if they are within a distance of $2\epsilon = 0.2$ within any target point in $\mathcal{Y}$. The solid points are selected as a random subset of the target point cloud, but we add a small random spatial offset to every solid query point that itself is uniformly sampled within a spherical ball of radius $\epsilon = 0.1$. This roughly corresponds to the spacing between target points on the objects and the floor in the dataset, encouraging the model to learn to "fill in the gap" between those points.

In the case of CARLA, to ensure that we maintain an ef-

fective learning signal, we describe several tricks that help guide supervision toward areas where it is deemed more important. In addition to random sampling from the target point cloud, 14% of solid points are explicitly sampled in dynamic regions of the scene, *i.e.* moving points that did not exist in a randomly selected other frame. The converse is also done for air points (*i.e.* regions that are now missing, but were present in another frame). At least 7% of sampled solid points focus exclusively on vehicles and pedestrians. We also oversample occluded vehicles and pedestrians (for up to 7% of all sampled solid points), by counting and comparing the number of points of every object seen by every view. Moreover, we ensure that objects that were never seen in the first place (*e.g.* pedestrians who remain behind a building or wall throughout the entire video) are not oversampled, to avoid confusing the model. Lastly, for class balancing, 14% of sampled solid points treat all semantic categories equally, *i.e.* we sample the same number of points from every class that is present in the scene.

For a fair and correct evaluation, there are no sampling tricks at test time, *i.e.* we apply uniform sampling within the cuboid of interest in a way that is agnostic of the ground truth. For GREATER, we sample $N = 2^{19}$ points per frame per scene, while for CARLA, $N = 2^{21}$.

A.3. Evaluation Metrics

The model is evaluated by the Chamfer Distance (CD) between the sampled prediction and the ground truth point cloud. Sometimes, the model fails to predict an object when it is fully occluded (*i.e.* a false negative), which may cause the output point cloud (filtered by the desired category and occlusion rate) to consist of zero points. Rare or "tiny" classes with a low number of target points per scene, *e.g.* traffic sign, may face similar issues. In that case, the CD would normally be undefined, but in order to ensure that the average metric accounts for this and is still affected, we substitute the prediction with a single point in the center of the scene: $(0, 0, 0)$ for GREATER, or $(20, 0, 0)$ for CARLA. Moreover, unless explicitly noted otherwise, we set $\sigma_T = 0.4$ for numerical evaluations (*i.e.* Tables 1 and 2) to reduce the likelihood of false negatives, and $\sigma_T = 0.5$ for visualization purposes (*i.e.* Figures 5 and 6).

B. Dataset Description

Both GREATER and CARLA are posed multiview RGB-D video datasets, with added instance segmentation and semantic segmentation annotations respectively. The camera views are illustrated in Figures 11 and 12.

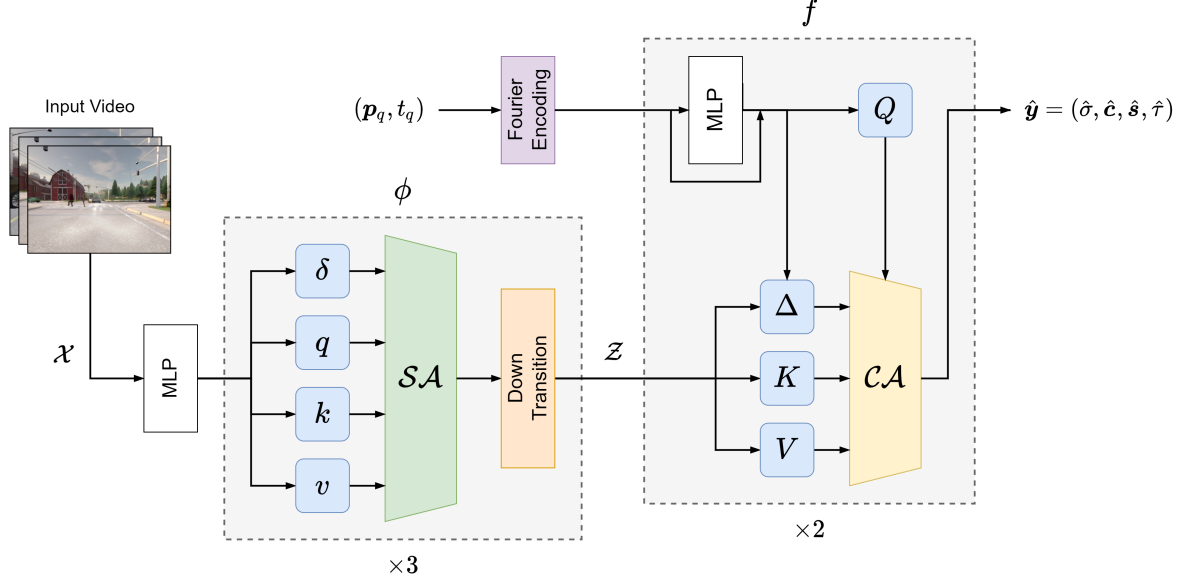


Figure 10. **Detailed Architecture** – We show the feature encoder ϕ and implicit representation f , with self-attention $\mathcal{SA}$ and cross-attention $\mathcal{CA}$ operations respectively. ϕ is a point transformer with four self-attention layers in total (the one after the last down transition is not shown), although the outputs of the last two down self-attention blocks are combined to form multi-scale features for $\mathcal{Z}$. The exact operation of the down transition modules is described in [69]. For the query points $(\mathbf{p}_q, t_q)$, we adopt the same Fourier encoding as [34]. All MLP blocks consist of two linear layers with a ReLU non-linearity in-between. The MLP blocks in the implicit representation f are residual, similar to [67].

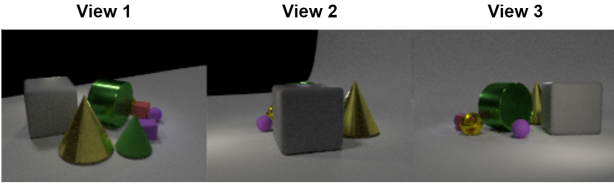


Figure 11. **GREATER Dataset Views** – Every GREATER scene has three non-moving RGB-D cameras at uniformly random azimuth angles, with the condition that all three are spaced by at least 45° away from each other. A random view is always selected to serve as input view, such that the other two views (along with the input view itself) serve as supervision during training.

B.1. GREATER

All videos are recorded with a virtual RGB-D camera with known intrinsics and extrinsics. The dataset is generated at 24 frames per second (FPS), and objects move and/or rotate in synchronized cycles that repeat every 32 to 42 frames. The model’s data loader subsamples temporally and uses 8 FPS, such that along with $T = 12$, roughly one full cycle is covered per clip. All objects are precisely twice as large as compared to CATER [18]. For every scene, the number of objects is selected uniformly at random between

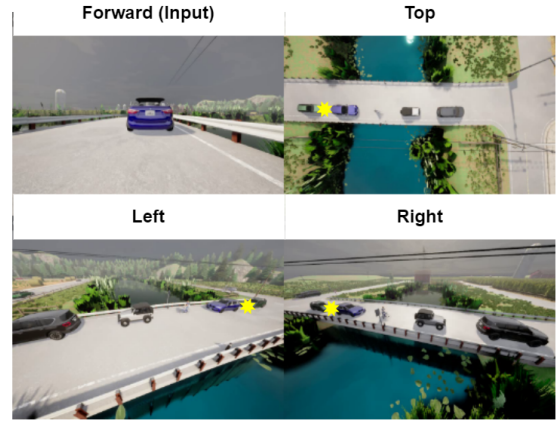


Figure 12. **CARLA Dataset Views** – Every CARLA scene has one RGB and LiDAR sensor attached to the front of the ego vehicle, which always records the input video. Three other supervisory views (in which the forward sensor position is marked with a yellow asterisk) operate only at training time.

8 and 12 (inclusive). The cameras are also chosen randomly per scene, but henceforth remain static (*i.e.* never move over time) over the duration of a single scene. With the spatially vertical axis denoted z , the 3D bounding box within

which both the input and predictions happen is $x \in [-5, 5]$, $y \in [-5, 5]$, $z \in [-1, 5]$.

B.2. CARLA

All videos are recorded with a pair of sensors with known intrinsics and extrinsics: one RGB camera, and one LiDAR sensor. The point cloud data generated by the latter sensor does not contain color information, so we use the RGB images to colorize the points. While the dataset is generated at 10 FPS, the model’s data loader subsamples temporally and uses 5 FPS. Both the LiDAR and camera horizontal fields of view are 120 degrees. However, the LiDAR’s spherical geometry is different from a camera’s projective geometry, implying that the LiDAR data is not directly aligned with the camera intrinsics. Therefore, in order to obtain a colorized point cloud per frame, we first project all LiDAR points onto the image and then map the pixel’s RGB values it was assigned to back to the 3D point. If a LiDAR point falls outside of the camera’s field of view, we mark it with a generic “color unavailable” constant, *i.e.* $(-1, -1, -1)$.

For every input clip passed to the model, since the vehicle pose over time is known, we correct all point clouds to a common reference frame. This reference frame is chosen to be the last input (and output) frame, such that the pair of sensors mounted to the ego vehicle is always at $(0, 0, 1)$ at time $t = T - 1$. With the forward axis denoted x , the sideways axis denoted y , and the vertical axis z , the input bounding box (*i.e.* containing all observed points) is $x \in [-14, 50]$, $y \in [-20, 20]$, $z \in [-1, 10]$, and the output bounding box (*i.e.* containing all predicted and ground truth points) is $x \in [0, 40]$, $y \in [-16, 16]$, $z \in [-1, 6.4]$ (in meters).

B.3. Clip Sampling

During training, for GREATER, we sample clips uniformly randomly. For CARLA however, most clips are relatively uninteresting, and we encourage learning about occlusions by performing biased clip sampling. Specifically, we construct a subset of starting frame indices where we know (as derived from semantic information over time in the LiDAR point clouds provided by the simulator) that occlusions are more likely to happen, and return a clip from this pool 40% of the time. To preemptively avoid overfitting, the data loader will never return the exact same clip twice over the entire duration of training.

During testing and evaluation, we deterministically sample a single clip within every video that has the most occlusions happening at once, counted over the number of objects. This implies that the test set for CARLA is significantly more challenging than the average driving situation (*i.e.* as compared to if we were to sample clips uniformly at random).

C. Qualitative Results

Figures 13 and 14 showcase *non-cherry-picked* visualizations made by our model as well as two important baselines. In terms of predictions, we depict our non-ablated 4D dynamic scene completion model, the ablation without time (which essentially becomes a 3D scene completion model without dynamics), and finally the PCN [68] baseline. Note that even though our adaptation of PCN can see temporal context and is given an advantage by borrowing features from the ground truth point cloud, the scenes in our dataset appear to be too complicated for PCN to learn effectively, especially in the case of CARLA.

Finally, please see our webpage at occlusions.cs.columbia.edu for more visualizations, as well as links to our datasets, source code, and models.

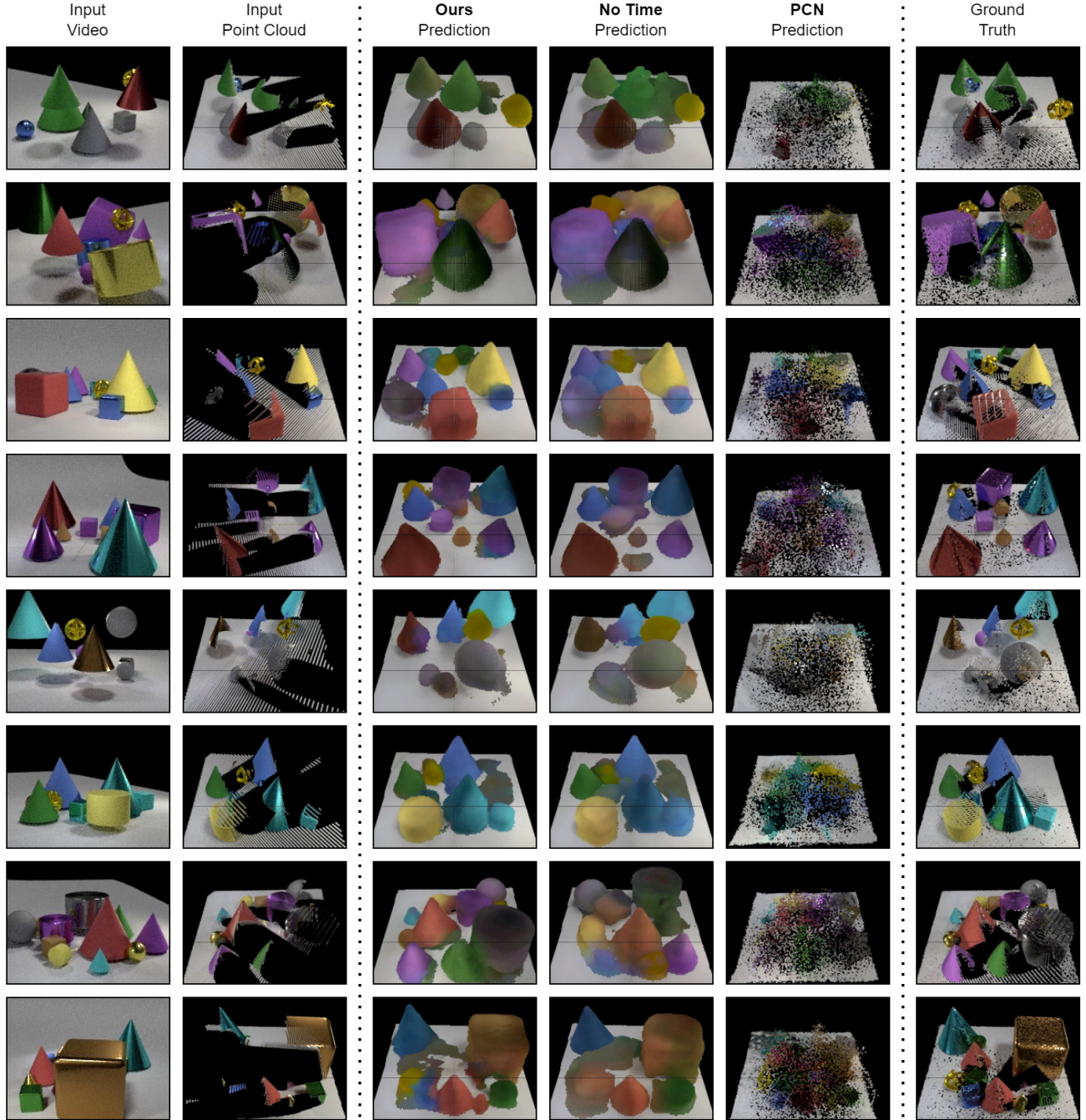


Figure 13. **Non-cherry-picked results for GREATER** – We show the last input frame, last output frames, and corresponding ground truth point cloud. Note that PCN does not have a mechanism built-in for predicting features (in this case, RGB color) associated with points, so we *copy colors from the nearest ground truth point* in order to boost the output legibility. Our model qualitatively has the best understanding of object permanence, and visually outperforms the baselines.

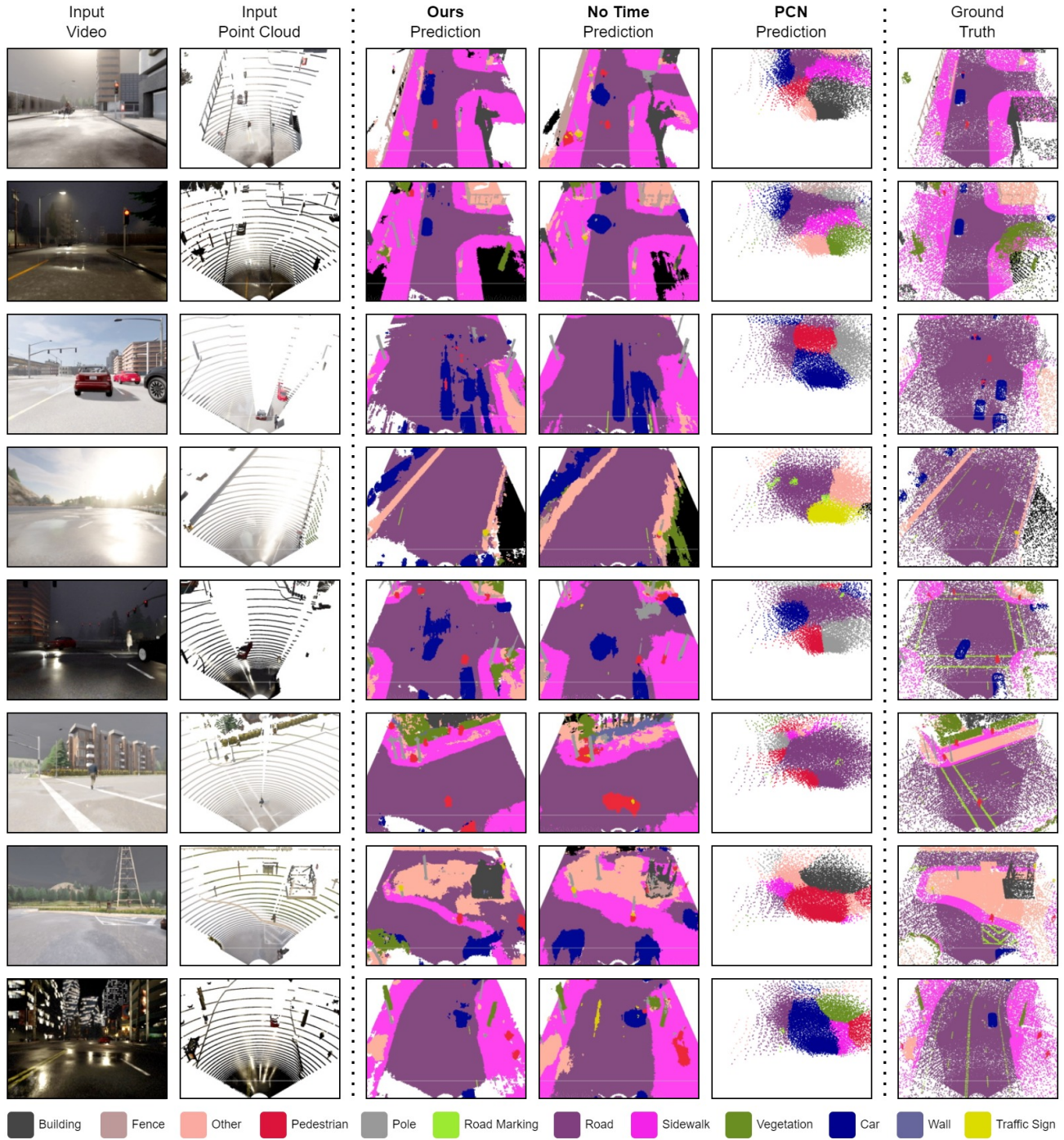


Figure 14. **Non-cherry-picked results for CARLA** – We show the last input frame, last output frames, and corresponding ground truth point cloud. Note that PCN does not have a mechanism built-in for predicting features (in this case, semantic category) associated with points, so we *copy category information from the nearest ground truth point* in order to boost the output legibility. Our model qualitatively has the best understanding of object permanence, and visually outperforms the baselines.