

Deep Stereo using Adaptive Thin Volume Representation with Uncertainty Awareness

Shuo Cheng^{1*} Zexiang Xu^{1*} Shilin Zhu¹

Zhuwen Li² Li Erran Li^{3,4} Ravi Ramamoorthi¹ Hao Su¹

¹University of California, San Diego ²Nuro Inc. ³Scale AI ⁴Columbia University

Abstract

We present *Uncertainty-aware Cascaded Stereo Network (UCS-Net)* for 3D reconstruction from multiple RGB images. Multi-view stereo (MVS) aims to reconstruct fine-grained scene geometry from multi-view images. Previous learning-based MVS methods estimate per-view depth using plane sweep volumes (PSVs) with a fixed depth hypothesis at each plane; this requires densely sampled planes for high accuracy, which is impractical for high-resolution depth because of limited memory. In contrast, we propose *adaptive thin volumes (ATVs)*; in an ATV, the depth hypothesis of each plane is spatially varying, which adapts to the uncertainties of previous per-pixel depth predictions. Our UCS-Net has three stages: the first stage processes a small PSV to predict low-resolution depth; two ATVs are then used in the following stages to refine the depth with higher resolution and higher accuracy. Our ATV consists of only a small number of planes with low memory and computation costs; yet, it efficiently partitions local depth ranges within learned small uncertainty intervals. We propose to use variance-based uncertainty estimates to adaptively construct ATVs; this differentiable process leads to reasonable and fine-grained spatial partitioning. Our multi-stage framework progressively sub-divides the vast scene space with increasing depth resolution and precision, which enables reconstruction with high completeness and accuracy in a coarse-to-fine fashion. We demonstrate that our method achieves superior performance compared with other learning-based MVS methods on various challenging datasets.

1. Introduction

Inferring 3D scene geometry from captured images is a core problem in computer vision and graphics with applications in 3D visualization, scene understanding, robotics and autonomous driving. Multi-view stereo (MVS) aims to reconstruct dense 3D representations from multiple images

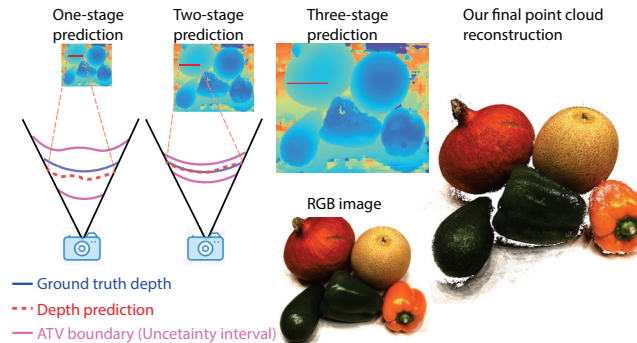


Figure 1: Our UCS-Net leverages adaptive thin volumes (ATVs) to progressively reconstruct a highly accurate high-resolution depth map through multiple stages. We show the input RGB image, depth predictions with increasing sizes from three stages, and our final point cloud reconstruction obtained by fusing multiple depth maps. We also illustrate a local slice (red) from our depth prediction with the corresponding ATV boundaries that reflect pixel-wise uncertainty intervals. Our ATVs become thinner after a stage with reduced uncertainty, which enables higher accuracy.

with calibrated cameras. Inspired by the success of deep convolutional neural networks (CNN), several learning-based MVS methods have been presented [23, 27, 54, 20, 47]; the most recent work leverages cost volumes in a learning pipeline [58, 21], and outperforms many traditional MVS methods [13].

At the core of the recent success on MVS [58, 21] is the application of 3D CNNs on plane sweep cost volumes to effectively infer multi-view correspondence. However, such 3D CNNs involve massive memory usage for depth estimation with high accuracy and completeness. In particular, for a large scene, high accuracy requires sampling a large number of sweeping planes and high completeness requires reconstructing high-resolution depth maps. In general, given limited memory, there is an undesired trade-off between accuracy (more planes) and completeness (more pixels) in previous work [58, 21].

* Equal contribution.

Our goal is to achieve *highly accurate and highly complete reconstruction* with *low memory and computation consumption* at the same time. To do so, we propose a novel learning-based uncertainty-aware multi-view stereo framework, which utilizes multiple small volumes, instead of a large standard plane sweep volume, to progressively regress high-quality depth in a coarse-to-fine fashion. A key in our method is that we propose novel adaptive thin volumes (ATVs, see Fig. 1) to achieve efficient spatial partitioning.

Specifically, we propose a novel cascaded network with three stages (see Fig. 2): each stage of the cascade predicts a depth map with a different size; each following stage constructs an ATV to refine the predicted depth from the previous stage with higher pixel resolution and finer depth partitioning. The first stage uses a small standard plane sweep volume with low image resolution and relatively sparse depth planes – 64 planes that are fewer than the number of planes (256 or 512) in previous work [58, 59]; the following two stages use ATVs with higher image resolutions and significantly fewer depth planes – only 32 and 8 planes. While consisting of a very small number of planes, our ATVs are constructed within *learned local depth ranges*, which enables *efficient and fine-grained spatial partitioning* for accurate and complete depth reconstruction.

This is made possible by the novel uncertainty-aware construction of an ATV. In particular, we leverage the variances of the predicted per-pixel depth probabilities, and infer the uncertainty intervals (as shown in Fig. 1) by calculating variance-based confidence intervals of the per-pixel probability distributions for the ATV construction. Specifically, we apply the previously predicted depth map as a central curved plane, and construct an ATV around the central plane within local per-pixel uncertainty intervals. In this way, we explicitly express the uncertainty of the depth prediction at one stage, and embed this knowledge into the input volume for the next stage.

Our variance-based uncertainty estimation is differentiable and we train our UCSNet from end to end with depth supervision for the predicted depths from all three stages. Our network can thus learn to optimize the estimated uncertainty intervals, to make sure that an ATV is constructed with proper depth coverage that is both large enough – to try to cover ground truth depth – and small enough – to enable accurate reconstruction for the following stages. Overall, our multi-stage framework can progressively sub-divide the local space at a finer scale in a reasonable way, which leads to high-quality depth reconstruction. We demonstrate that our novel UCS-Net outperforms the state-of-the-art learning-based MVS methods on various datasets.

2. Related Work

Multi-view stereo is a long-studied vision problem with many traditional approaches [44, 39, 33, 32, 26, 10, 8, 13,

43]. Our learning-based framework leverages the novel spatial representation, ATV to reconstruct high-quality depth for fine-grain scene reconstruction. In this work, we mainly discuss spatial representation for 3D reconstruction and deep learning based multi-view stereo.

Spatial Representation for 3D Reconstruction. Existing methods can be categorized based on learned 3D representations. Volumetric based approaches partition the space into a regular 3D volume with millions of small voxels [23, 27, 54, 55, 60, 40], and the network predicts if a voxel is on the surface or not. Ray tracing can be incorporated into this voxelized structure [49, 38, 50]. The main drawback of these methods is computation and memory inefficiency, given that most voxels are not on the surface. Researchers have also tried to reconstruct point clouds [22, 13, 35, 52, 36, 2], however the high dimensionality of a point cloud often results in noisy outliers since a point cloud does not efficiently encode connectivity between points. Some recent works utilize single or multiple images to reconstruct a point cloud given strong shape priors [11, 22, 36], which cannot be directly extended to large-scale scene reconstruction. Recent work also tried to directly reconstruct surface meshes [34, 25, 53, 19, 46, 28], deformable shapes [24, 25], and some learned implicit distance functions [7, 41, 37, 6]. These reconstructed surfaces often look smoother than point-cloud-based approaches, but often lack high-frequency details. A depth map represents dense 3D information that is perfectly aligned with a reference view; depth reconstruction has been demonstrated in many previous works on reconstruction with both single view [9, 51, 16, 17, 62] and multiple views [4, 48, 18, 14, 43, 57, 43]. Some of them leverage normal information as well [14, 15]. In this paper, we present ATV, a novel spatial representation for depth estimation; we use two ATVs to progressively partition local space, which is the key to achieve coarse-to-fine reconstruction.

Deep Multi-View Stereo (MVS). The traditional MVS pipeline mainly relies on photo-consistency constraints to infer the underlying 3D geometry, but usually performs poorly on texture-less or occluded areas, or under complex lighting environments. To overcome such limitations, many deep learning-based MVS methods have emerged in the last two years, including regression-based approaches [58, 21], classification-based approaches [20] and approaches based on recurrent- or iterative- style architectures [59, 61, 5] and many other approaches [30, 38, 3, 45]. Most of these methods build a single cost volume with uniformly sampled depth hypotheses by projecting 2D image features into 3D space, and then use a stack of either 2D or 3D CNNs to infer the final depth [58, 12, 56]. However, a single cost volume often requires a large number of depth planes to achieve enough reconstruction accuracy, and it is difficult to reconstruct high-resolution depth, limited by the memory bottle-

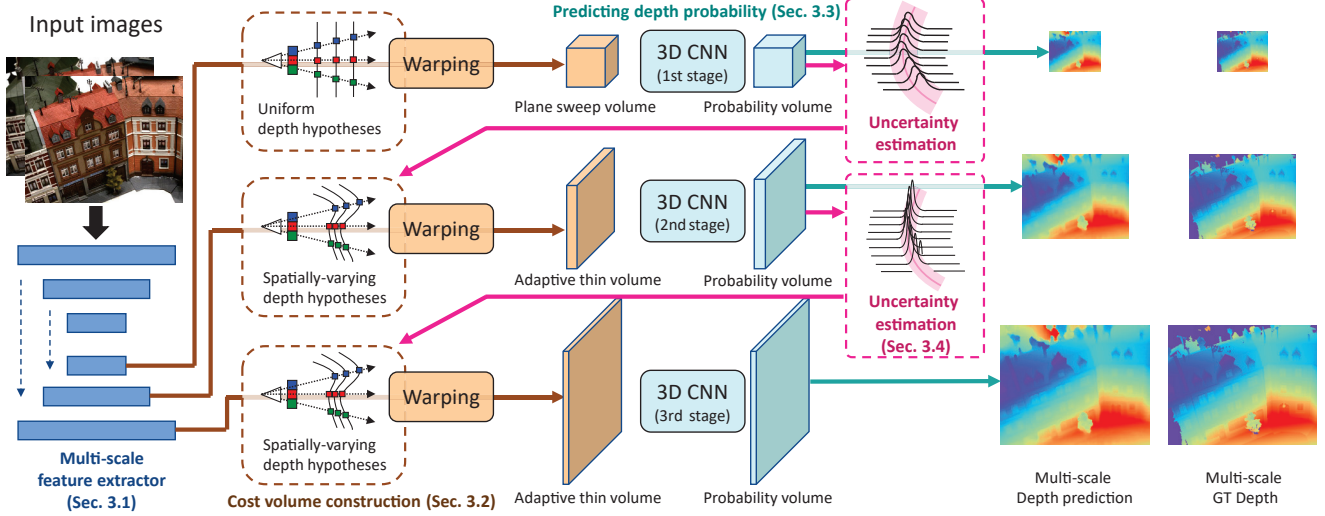


Figure 2: Overview of our UCS-Net. Our UCS-Net leverages multi-scale cost volumes to achieve coarse-to-fine depth prediction with three cascade stages. The cost volumes are constructed using multi-scale deep image features from a multi-scale feature extractor. The last two stages utilize the uncertainty of the previous depth prediction to build adaptive thin volumes (ATVs) for depth reconstruction at a finer scale. We mark different parts of the network in different colors. Please refer to Sec 3 and the corresponding subsections for more details.

neck. R-MVSNet [59] leverages recurrent networks to sequentially build a cost volume with a high depth-wise sampling rate (512 planes). In contrast, we apply an adaptive sampling strategy with ATVs, which enables more efficient spatial partitioning with a higher depth-wise sampling rate using fewer depth planes (104 planes in total, see Tab. 3), and our method achieves significantly better reconstruction than R-MVSNet (see Tab. 1 and Tab. 2). On the other hand, Point-MVSNet [5] densifies a coarse reconstruction within a predefined local spatial range for better reconstruction with learning-based refinement. We propose to refine depth in a learned local space with adaptive thin volumes to obtain accurate high-resolution depth, which leads to better reconstruction than Point-MVSNet and other state-of-the-art methods (see Tab. 1 and Tab. 2).

3. Method

Some recent works aim to improve learning-based MVS methods. Recurrent networks [59] have been utilized to achieve fine depth-wise partitioning for high accuracy; a PointNet-based method [5] is also presented to densify the reconstruction for high completeness. Our goal is to reconstruct high-quality 3D geometry with both high accuracy and high completeness. To this end, we propose a novel uncertainty-aware cascaded network (UCS-Net) to reconstruct highly accurate per-view depth with high resolution.

Given a reference image I_1 and $N - 1$ source images $\{I_i\}_{i=2}^N$, our UCS-Net progressively regresses a fine-grained depth map at the same resolution as the refer-

ence image. We show the architecture of the UCS-Net in Fig. 2. Our UCS-Net first leverages a 2D CNN to extract multi-scale deep image features at three resolutions (Sec. 3.1). Our depth prediction is achieved through three stages, which leverage multi-scale image features to predict multi-resolution depth maps. In these stages, we construct multi-scale cost volumes (Sec. 3.2), where each volume is a plane sweep volume or an adaptive thin volume (ATV). We then apply 3D CNNs to process the cost volumes to predict per-pixel depth probability distributions, and a depth map is reconstructed from the expectations of the distributions (Sec. 3.3). To achieve efficient spatial partitioning, we utilize the uncertainty of the depth prediction to construct ATVs as cost volumes for the last two stages (Sec. 3.4). Our multi-stage network effectively reconstructs depth in a coarse-to-fine fashion (Sec. 3.5).

3.1. Multi-scale feature extractor

Previous methods use downsampling layers [58, 59] or a UNet [56] to extract deep features and build a plane sweep volume at a single resolution. To reconstruct high-resolution depth, we introduce a multi-scale feature extractor, which enables constructing multiple cost volumes at different scales for multi-resolution depth prediction. As schematically shown in Fig. 2, our feature extractor is a small 2D UNet [42], which has an encoder and a decoder with skip connections. The encoder consists of a set of convolutional layers followed by BN (batch normalization) and ReLu activation layers; we use stride = 2 convolutions to

downsample the original image size twice. The decoder upsamples the feature maps, convolves the upsampled features and the concatenated features from skip links, and also applies BN and Relu layers. Given each input image $\mathbf{I}_i$, the feature extractor provides three scale feature maps, $F_{i,1}$, $F_{i,2}$, $F_{i,3}$, from the decoder for the following cost volume construction. We represent the original image size as $W \times H$, where W and H denote the image width and height; correspondingly, $F_{i,1}$, $F_{i,2}$ and $F_{i,3}$ have resolutions of $\frac{W}{4} \times \frac{H}{4}$, $\frac{W}{2} \times \frac{H}{2}$ and $W \times H$, and their numbers of channels are 32, 16 and 8 respectively. Our multi-scale feature extractor allows for the high-resolution features to properly incorporate the information at lower resolutions through the learned upsampling process; thus in the multi-stage depth prediction, each stage is aware of the meaningful feature knowledge used in previous stages, which leads to reasonable high-frequency feature extraction.

3.2. Cost volume construction

We construct multiple cost volumes at multiple scales by warping the extracted feature maps, $F_{i,1}$, $F_{i,2}$, $F_{i,3}$ from source views to a reference view. Similar to previous work, this process is achieved through differentiable unprojection and projection. In particular, given camera intrinsic and extrinsic matrices $\{K_i, T_i\}$ for each view i , the 4×4 warping matrix at depth d at the reference view is given by:

$$H_i(d) = K_i T_i T_1^{-1} K_1^{-1}. \quad (1)$$

In particular, when warping to a pixel in the reference image $\mathbf{I}_1$ at location (x, y) and depth d , $H_i(d)$ multiplies the homogeneous vector $(xd, yd, d, 1)$ to find its corresponding pixel location in each $\mathbf{I}_i$ in homogeneous coordinates.

Each cost volume consists of multiple planes; we use $\mathbf{L}_{k,j}$ to denote the depth hypothesis of the j th plane at the k th stage, and $\mathbf{L}_{k,j}(x)$ represents its value at pixel x . At stage k , once we warp per-view feature maps $F_{i,k}$ at all depth planes with corresponding hypotheses $\mathbf{L}_{k,j}$, we calculate the variance of the warped feature maps across views at each plane to construct a cost volume. We use D_k to represent the number of planes for stage k . For the first stage, we build a standard plane sweep volume, whose depth hypotheses are of constant values, i.e. $\mathbf{L}_{1,j}(x) = d_j$. We uniformly sample $\{d_j\}_{j=1}^{D_1}$ from a pre-defined depth interval $[d_{min}, d_{max}]$ to construct the volume, in which each plane is constructed using $H_i(d_j)$ to warp multi-view images. For the second and third stages, we build novel adaptive thin volumes, whose depth hypotheses have spatially-varying depth values according to pixel-wise uncertainty estimates of the previous depth prediction. In this case, we calculate per-pixel per-plane warping matrices by setting $d = \mathbf{L}_{k,j}(x)$ in Eqn. 1 to warp images and construct cost volumes. Please refer to Sec. 3.4 for uncertainty estimation.

3.3. Depth prediction and probability distribution

At each stage, we apply a 3D CNN to process the cost volume, infer multi-view correspondence and predict depth probability distributions. In particular, we use a 3D UNet similar to [58], which has multiple downsampling and upsampling 3D convolutional layers to reason about scene geometry at multiple scales. We apply depth-wise softmax at the end of the 3D CNNs to predict per-pixel depth probabilities. Our three stages use the same network architecture without sharing weights, so that each stage learns to process its information at a different scale. Please refer to the supplemental material for details of our 3D CNN architecture.

The 3D CNN at each stage predicts a depth probability volume that consists of D_k depth probability maps $\mathbf{P}_{k,j}$ associated with the depth hypotheses $\mathbf{L}_{k,j}$. $\mathbf{P}_{k,j}$ expresses per-pixel depth probability distributions, where $\mathbf{P}_{k,j}(x)$ represents how probable the depth at pixel x is $\mathbf{L}_{k,j}(x)$. A depth map $\hat{\mathbf{L}}_k$ at stage k is reconstructed by weighted sum:

$$\hat{\mathbf{L}}_k(x) = \sum_{j=1}^{D_k} \mathbf{L}_{k,j}(x) \cdot \mathbf{P}_{k,j}(x). \quad (2)$$

3.4. Uncertainty estimation and ATV

The key for our framework is to progressively sub-partition the local space and refine the depth prediction with increasing resolution and accuracy. To do so, we construct novel ATVs for the last two stages, which have curved sweeping planes with spatially-varying depth hypotheses (as illustrated in Fig. 1 and Fig. 2), based on uncertainty inference of the predicted depth in its previous stage.

Given a set of depth probability maps, previous work only utilizes the expectation of the per-pixel distributions (using Eqn. (2)) to determine an estimated depth map. For the first time, we leverage the variance of the distribution for uncertainty estimation, and construct ATVs using the uncertainty. In particular, the variance $\hat{\mathbf{V}}_k(x)$ of the probability distribution at pixel x and stage k is calculated as:

$$\hat{\mathbf{V}}_k(x) = \sum_{j=1}^{D_k} \mathbf{P}_{k,j}(x) \cdot (\mathbf{L}_{k,j}(x) - \hat{\mathbf{L}}_k(x))^2, \quad (3)$$

and the corresponding standard deviation is $\hat{\sigma}_k(x) = \sqrt{\hat{\mathbf{V}}_k(x)}$. Given the depth prediction $\hat{\mathbf{L}}_k(x)$ and its variance $\hat{\sigma}_k(x)^2$ at pixel x , we propose to use a variance-based confidence interval to measure the uncertainty of the prediction:

$$\mathbf{C}_k(x) = [\hat{\mathbf{L}}_k(x) - \lambda \hat{\sigma}_k(x), \hat{\mathbf{L}}_k(x) + \lambda \hat{\sigma}_k(x)], \quad (4)$$

where λ is a scalar parameter that determines how large the confidence interval is. For each pixel x , we uniformly sample D_{k+1} depth values from $\mathbf{C}_k(x)$ of the k th stage, to get its depth values $\mathbf{L}_{k+1,1}(x), \mathbf{L}_{k+1,2}(x), \dots, \mathbf{L}_{k+1,D_{k+1}}(x)$

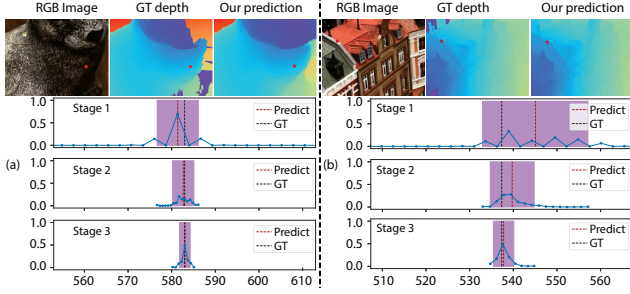


Figure 3: We illustrate detailed depth and uncertainty estimation of two examples. On the top, we show the RGB image crops, predicted depth and ground truth depth. On the bottom, we show the details of two pixels (red points in the images) with predicted depth probabilities (connected blue dots), depth prediction (red dash line), the ground truth depth (black dash line) and uncertainty intervals (purple) in the three stages.

of the depth planes for stage $(k + 1)$. In this way, we construct D_{k+1} spatially-varying depth hypotheses $\mathbf{L}_{k+1,j}$, which form the ATV for stage $(k + 1)$.

The estimated $\mathbf{C}_k(x)$ expresses the uncertainty interval of the prediction $\hat{\mathbf{L}}_k(x)$, which determines the physical thickness of an ATV at each pixel. In Fig. 3, we show two actual examples with two pixels and their estimated uncertainty intervals $\mathbf{C}_k(x)$ around the predictions (red dash line). The $\mathbf{C}_k$ essentially depicts a probabilistic local space around the ground truth surface, and the ground truth depth is located in the uncertainty interval with a very high confidence. Note that, our variance-based uncertainty estimation is differentiable, which enables our UCS-Net to learn to adjust the probability prediction at each stage to achieve optimized intervals and corresponding ATVs for following stages in an end-to-end training process. As a result, the spatially varying depth hypotheses in ATVs naturally adapt to the uncertainty of depth predictions, which leads to highly efficient spatial partitioning.

3.5. Coarse-to-fine prediction

Our UCS-Net leverages three stages to reconstruct depth at multiple scales from coarse to fine, which generally supports different numbers (D_k) of planes in each stage. In practice, we use $D_1 = 64$, $D_2 = 32$ and $D_3 = 8$ to construct a plane sweep volume and two ATVs with sizes of $\frac{W}{4} \times \frac{H}{4} \times 64$, $\frac{W}{2} \times \frac{H}{2} \times 32$ and $H \times W \times 8$ to estimate depth at corresponding resolutions. While our two ATVs have small numbers (32 and 8) of depth planes, they in fact partition local depth ranges at finer scales than the first stage volume; this is achieved by our novel uncertainty-aware volume construction process which adaptively controls local depth intervals. This efficient usage of a small

Method	Acc.	Comp.	Overall
Camp [4]	0.835	0.554	0.695
Furu [13]	0.613	0.941	0.777
Tola [48]	0.342	1.190	0.766
Gipuma [14]	0.283	0.873	0.578
SurfaceNet [46]	0.450	1.040	0.745
MVSNet [58]	0.396	0.527	0.462
R-MVSNet [59]	0.383	0.452	0.417
Point-MVSNet [5]	0.342	0.411	<u>0.376</u>
Our 1st stage	0.548	0.529	0.539
Our 2nd stage	0.401	<u>0.397</u>	0.399
Our full model	<u>0.338</u>	0.349	0.344

Table 1: Quantitative results of accuracy, completeness and overall on the DTU testing set. Numbers represent distances in millimeters and smaller means better.

number of depth planes enables the last two stages to deal with higher pixel-wise resolutions given the limited memory, which makes fine-grained depth reconstruction possible. Our novel ATV effectively expresses the locality and uncertainty in the depth prediction, which enables state-of-the-art depth reconstruction results with high accuracy and high completeness through a coarse-to-fine framework.

3.6. Training details

Training set. We train our network on the DTU dataset [1]. We split the dataset into training, validate and testing set, and create ground truth depth similar to [58]. In particular, we apply Poisson reconstruction [29] on the point clouds in DTU, and render the surface at the captured views with three resolutions, $\frac{W}{4} \times \frac{H}{4}$, $\frac{W}{2} \times \frac{H}{2}$ and the original $W \times H$. In particular, we use $W \times H = 640 \times 512$ for training.

Loss function. Our UCS-Net predicts depth at three resolutions; we apply $L1$ loss on depth prediction at each resolution with the rendered ground truth at the same resolution. Our final loss is the combination of the three $L1$ losses.

Training policy. We train our full three-stage network from end to end for 60 epochs. We use Adam optimizer with an initial learning rate of 0.0016. We use 8 NVIDIA GTX 1080Ti GPUs to train the network with a batch size of 16 (mini-batch size of 2 per GPU).

4. Experiments

We now evaluate our UCS-Net. We do benchmarking on the DTU and Tanks and Temple datasets. We then justify the effectiveness of the designs of our network, in terms of uncertainty estimation and multi-stage prediction.

Evaluation on the DTU dataset [1]. We evaluate our method on the DTU testing set. To reconstruct the final point cloud, we follow [14] to fuse the depth from multiple views; we use this fusion method for all our exper-

Method	Mean	Family	Francis	Horse	Lighthouse M60	Panther	Playground Train		
MVSNet[58]	43.48	55.99	28.55	25.07	<u>50.79</u>	<u>53.96</u>	<u>50.86</u>	47.90	34.69
R-MVSNet[59]	48.40	69.96	46.65	32.59	42.95	51.88	48.80	52.00	42.38
Dense-R-MVSNet[59]	<u>50.55</u>	<u>73.01</u>	54.46	43.42	43.88	46.80	46.69	50.87	<u>45.25</u>
Point-MVSNet[5]	48.27	61.79	41.15	34.20	<u>50.79</u>	51.97	50.85	<u>52.38</u>	43.06
Our full model	54.83	76.09	<u>53.16</u>	<u>43.03</u>	54.00	55.60	51.49	57.38	47.89

Table 2: Quantitative results of F-scores (higher means better) on Tanks and Temples.

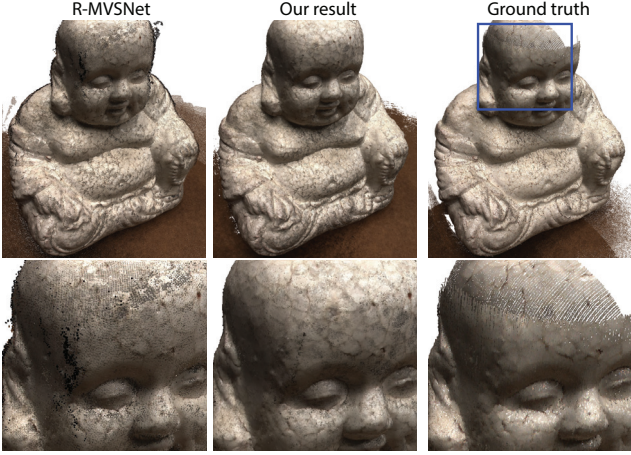


Figure 4: Comparisons with R-MVSNet on an example in the DTU dataset. We show rendered images of the point clouds of our method, R-MVSNet and the ground truth. In this example, the ground truth from scanning is incomplete. We also show insets for detailed comparisons marked as a blue box in the ground truth. Note that our result is smoother and has fewer outliers than R-MVSNet’s result.

iments. For fair comparisons, we use the same view selection, image size and initial depth range as in [58] with $N = 5$, $W = 1600$, $H = 1184$, $d_{\min} = 425mm$ and $d_{\max} = 933.8mm$; similar settings are also used in other learning-based MVS methods [5, 59]. We use a NVIDIA GTX 1080 Ti GPU to run the evaluation.

We compare the accuracy and the completeness of the final reconstructions using the distance metric in [1]. We compare against both traditional methods and learning-based methods, and the average quantitative results are shown in Tab. 1. While Gipuma [14] (a traditional method) achieves the best accuracy among all methods, our method has significantly better completeness and overall scores. Besides, our method outperforms all state-of-the-art baseline methods in terms of both accuracy and completeness. Note that with the same input, MVSNet and R-MVSNet predict depth maps with a size of only $\frac{W}{4} \times \frac{H}{4}$; our final depth maps are estimated at the original image size, which are of much higher resolution and lead to significantly better completeness. Meanwhile, such high completeness is

obtained without losing accuracy; our accuracy is also significantly better thanks to our uncertainty-aware progressive reconstruction. Point-MVSNet [5] densifies low-resolution depth within a predefined local depth range, which also reconstructs depth at the original image resolution; in contrast, our UCS-Net leverages learned adaptive local depth ranges and achieves better accuracy and completeness.

We also show results from our intermediate low-resolution depth of the first and the second stages in Tab. 1. Note that, because of sparser depth planes, our first-stage results (64 planes) are worse than MVSNet (256 planes) and R-MVSNet (512 planes) that reconstruct depth at the same low resolution. Nevertheless, our novel uncertainty-aware network introduces highly efficient spatial partitioning with ATVs in the following stages, so that our intermediate second-stage reconstruction is already much better than the two previous methods, and our third stage further improves the quality and achieves the best reconstruction.

We show qualitative comparisons between our method and R-MVSNet [59] in Fig. 4, in which we use the released point cloud reconstruction on R-MVSNet’s website for the comparison. While both methods achieve comparable completeness in this example, it is very hard for R-MVSNet to achieve high accuracy at the same time, which introduces obvious outliers and noise on the surface. In contrast, our method is able to obtain high completeness and high accuracy simultaneously as reflected by the smooth complete geometry in the image.

Evaluation on Tanks and Temple dataset [31]. We now evaluate the generalization of our model by testing our network trained with the DTU dataset on complex outdoor scenes in the Tanks and Temple intermediate dataset. We use $N = 5$ and $W \times H = 1920 \times 1056$ for this experiment. Our method outperforms most published methods, and to the best of our knowledge, when comparing with all published learning-based methods, we achieve the best average F-score (54.83) as shown in Tab. 2. In particular, our method obtains higher F-scores than MVSNet [58] and Point-MVSNet [5] in all nine testing scenes. Dense-R-MVSNet leverages a well-designed post-processing method and achieves slightly better performance than ours on two of the scenes, whereas our work is focused on high-quality per-view depth reconstruction, and we use a traditional fu-

	Ratio	Interval	D_k	Unit
PSV	100%	508.8mm	64	7.95mm
1st ATV	94.72%	13.88mm	32	0.43mm
2st ATV	85.22%	3.83mm	8	0.48mm

Table 3: Evaluation of uncertainty estimation. The PSV is the first-stage plane sweep volume; the 1st ATV is constructed after the first stage and used in the second stage; the 2nd ATV is used in the third stage. We show the percentages of uncertainty intervals that cover the ground truth depth. We also show the average length of the intervals, the number of depth planes and the unit sampling distance.

sion technique for post-processing. Nonetheless, thanks to our high-quality depth, our method still outperforms DenseR-MVSNet on most of the testing scenes and achieves the best overall performance.

Evaluation of uncertainty estimation. One key design of our UCS-Net is leveraging differentiable uncertainty estimation for the ATV construction. We now evaluate our uncertainty estimation on the DTU validate set. In Tab. 3, we show the average length of our estimated uncertainty intervals, the corresponding average sampling distances between planes, and the ratio of the pixels whose estimated uncertainty intervals cover the ground truth depth in the ATVs; we also show the corresponding values of the standard plane sweep volume (PSV) used in the first stage, which has an interval length of $d_{\max} - d_{\min} = 508.8\text{mm}$ and covers the ground truth depth with certainty.

We can see that our method is able to construct efficient ATVs that cover very local depth ranges. The first ATV significantly reduces the initial depth range from 508.8mm to only 13.88mm in average, and the second ATV further reduces it to only 3.83mm. Our ATV enables efficient depth sampling in an adaptive way, and obtains about 0.48mm sampling distance with only 32 or 8 depth planes. Note that, MVSNet and R-MVSNet sample the same large depth range (508.8mm) in a uniform way with a large number of planes (256 and 512); yet, the uniform sampling merely obtains volumes with sampling distances of 1.99mm and 0.99mm along depth. In contrast, our UCS-Net achieves a higher actual depth-wise sampling rate with a small number of planes; this allows for the focus of the cost volumes to be changed from sampling the depth to sampling the image plane with dense pixels in ATVs given the limited memory, which enables high-resolution depth reconstruction.

Besides, our adaptive thin volumes achieve high ratios (94.72% and 85.22%) of covering the ground truth depth in the validate set, as shown in Tab. 3; this justifies that our estimated uncertainty intervals are of high confidence. Our variance-based uncertainty estimation is equivalent to ap-

Stage	Scale	Size	Acc.	Comp.	Overall
1	$\times 1$	400x296	0.548	0.529	0.539
1	$\times 2$	800x592	0.411	0.535	0.473
2	$\times 1$	800x592	0.401	0.397	0.399
2	$\times 2$	1600x1184	0.342	0.386	0.364
3	$\times 1$	1600x1184	0.338	0.349	0.344

Table 4: Ablation study on the DTU testing set with different stages and upsampling scales (a scale of 1 represents the original result at the stage). The quantitative results represent average distances in mm (lower is better).

proximating a depth probability distribution as a Gaussian distribution and then computing its confidence interval with a specified scale on its standard deviation as in Eqn. 4.

We note that our variance-based uncertainty estimation is not only valid for single-mode Gaussian-like distributions as in Fig. 3.a, but also valid for many multi-mode cases as in Fig. 3.b, which shows a challenging example near object boundary. In Fig. 3.b, the predicted first-stage depth distribution has multiple modes; yet, it correspondingly has large variance and a large enough uncertainty interval. Our network predicts reasonable uncertainty intervals that are able to cover the ground truth depth in most cases, which allows for increasingly accurate reconstruction in the following stages at finer local spatial scales. This is made possible by the differentiable uncertainty estimation and the end-to-end training process, from which the network learns to control per-stage probability estimation to obtain proper uncertainty intervals for ATV construction. Because of this, we observe that our network is not very sensitive to different λ , and learns to predict similar uncertainty. Our uncertainty-aware volume construction process enables highly efficient spatial partitioning, which further allows for the final reconstruction to be of high accuracy and high completeness.

Evaluation of multi-stage depth prediction. We have quantitatively demonstrated that our multi-stage framework reconstructs scene geometry with increasing accuracy and completeness in every stage (see Fig. 1). We now further evaluate our network and do ablation studies about different stages on the DTU testing set with detailed quantitative and qualitative comparisons. We compare with naive upsampling to justify the effectiveness of our uncertainty-aware coarse-to-fine framework. In particular, we compare the results from our full model and the results from the first two stages with naive bilinear upsampling using a scale of 2 (for both height and width) in Tab. 4. We can see that upsampling does improve the reconstruction, which benefits from denser geometry and using our high-quality low-resolution results as input. However, the improvement made by naive upsampling is very limited, which is much lower than our improvement from our ATV-based upsampling. Our UCS-Net makes use of the ATV – a learned local spatial repre-

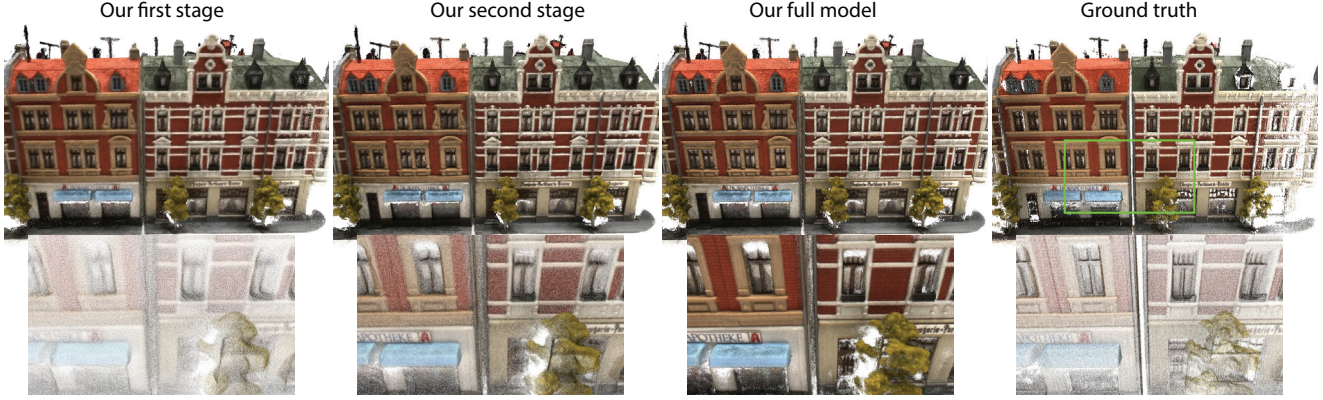


Figure 5: Qualitative comparisons between multi-stage point clouds and the ground truth point cloud on a scene in the DTU validate set. We show zoom-out (top) and zoom-in (bottom) rendered point clouds; the corresponding zoom-in region is marked in the ground truth as a green box. Our UCS-Net achieves increasingly dense and accurate reconstruction through the multiple stages. Note that, the ground truth point cloud is obtained by scanning, which is even of lower quality than our reconstructions in this example.

Method	Running time (s)	Memory (MB)	Input size	Prediction size
One stage	0.065	1309		160x120
Two stages	0.114	1607	640x480	320x240
Our full model	0.257	1647		640x480
MVSNet [58]	1.049	4511	640x480	160x120
R-MVSNet [59]	1.421	4261	640x480	160x120

Table 5: Performance comparisons. We show the running time and memory of our method by running the first stage, the first two stages and our full model.

sensation that is constructed in an uncertainty-aware way – to reasonably densify the map with a significant increase of both completeness and accuracy at the same time.

Figure 5 shows qualitative comparisons between our reconstructed point clouds and the ground truth point cloud. Our UCS-Net is able to effectively refine and densify the reconstruction through multiple stages. Note that, our MVS-based reconstruction is even more complete than the ground truth point cloud that is obtained by scanning, which shows the high quality of our reconstruction.

Comparing runtime performance. We now evaluate the timing and memory usage of our method. We run our model on the DTU validate set with an input image resolution of $W \times H = 640 \times 480$; We compare performance with MVSNet and R-MVSNet with 256 depth planes using the same inputs. Table 5 shows the performance comparisons including running time and memory. Note that, our full model is the only one that reconstructs the depth at the original image resolution that is much higher than the comparison methods. However, this hasn’t introduced any higher computation or memory consumption. In fact, our method requires signif-

icantly less memory and shorter running time, which are only about a quarter of the memory and time used in other methods. This demonstrates the benefits of our coarse-to-fine framework with fewer depth planes (104 in total), in terms of system resource usage. Our UCS-Net with ATVs achieves high-quality reconstruction with high computation and memory efficiency.

5. Conclusion

In this paper, we present a novel deep learning-based approach for multi-view stereo. We propose the novel uncertainty-aware cascaded stereo network (UCS-Net), which utilizes the adaptive thin volume (ATV), a novel spatial representation. For the first time, we make use of the uncertainty of the prediction in a learning-based MVS system. Specifically, we leverage variance-based uncertainty intervals at one cascade stage to construct an ATV for its following stage. The ATVs are able to progressively sub-partition the local space at a finer scale, and ensure that the smaller volume still surrounds the actual surface with a high probability. Our novel UCS-Net achieves highly accurate and highly complete scene reconstruction in a coarse-to-fine fashion. We compare our method with various state-of-the-art benchmarks; we demonstrate that our method is able to achieve the qualitatively and quantitatively best performance with high computation- and memory- efficiency. Our novel UCS-Net takes a step towards making the learning-based MVS method more reliable and efficient.

Acknowledgements This work was funded in part by Kuaishou Technology, NSF grant IIS-1764078, NSF grant 1703957, the Ronald L. Graham chair and the UC San Diego Center for Visual Computing.

References

- [1] Henrik Aanæs, Rasmus Ramsbøl Jensen, George Vogiatzis, Engin Tola, and Anders Bjarholm Dahl. Large-scale data for multiple-view stereopsis. *International Journal of Computer Vision*, 120(2):153–168, 2016. 5, 6
- [2] Panos Achlioptas, Olga Diamanti, Ioannis Mitliagkas, and Leonidas Guibas. Learning representations and generative models for 3d point clouds. In *International Conference on Machine Learning*, pages 40–49, 2018. 2
- [3] Konstantinos Batsos, Changjiang Cai, and Philippos Mordohai. CbmV: A coalesced bidirectional matching volume for disparity estimation. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 2060–2069, 2018. 2
- [4] Neill DF Campbell, George Vogiatzis, Carlos Hernández, and Roberto Cipolla. Using multiple hypotheses to improve depth-maps for multi-view stereo. In *European Conference on Computer Vision*, pages 766–779. Springer, 2008. 2, 5
- [5] Rui Chen, Songfang Han, Jing Xu, and Hao Su. Point-based multi-view stereo network. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 1538–1547, 2019. 2, 3, 5, 6
- [6] Zhiqin Chen and Hao Zhang. Learning implicit fields for generative shape modeling. *Proceedings of IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2019. 2
- [7] Angela Dai, Charles Ruizhongtai Qi, and Matthias Nießner. Shape completion using 3d-encoder-predictor cnns and shape synthesis. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 5868–5877, 2017. 2
- [8] Jeremy S De Bonet and Paul Viola. Poxels: Probabilistic voxelized volume reconstruction. In *Proceedings of International Conference on Computer Vision (ICCV)*, pages 418–425, 1999. 2
- [9] David Eigen and Rob Fergus. Predicting depth, surface normals and semantic labels with a common multi-scale convolutional architecture. In *Proceedings of the IEEE international conference on computer vision*, pages 2650–2658, 2015. 2
- [10] Carlos Hernández Esteban and Francis Schmitt. Silhouette and stereo fusion for 3d object modeling. *Computer Vision and Image Understanding*, 96(3):367–392, 2004. 2
- [11] Haoqiang Fan, Hao Su, and Leonidas J Guibas. A point set generation network for 3d object reconstruction from a single image. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 605–613, 2017. 2
- [12] John Flynn, Ivan Neulander, James Philbin, and Noah Snavely. Deepstereo: Learning to predict new views from the world’s imagery. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 5515–5524, 2016. 2
- [13] Yasutaka Furukawa and Jean Ponce. Accurate, dense, and robust multiview stereopsis. *IEEE transactions on pattern analysis and machine intelligence*, 32(8):1362–1376, 2010. 1, 2, 5
- [14] Silvano Galliani, Katrin Lasinger, and Konrad Schindler. Massively parallel multiview stereopsis by surface normal diffusion. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 873–881, 2015. 2, 5, 6
- [15] Silvano Galliani and Konrad Schindler. Just look at the image: viewpoint-specific surface normal prediction for improved multi-view reconstruction. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 5479–5487, 2016. 2
- [16] Ravi Garg, Vijay Kumar BG, Gustavo Carneiro, and Ian Reid. Unsupervised cnn for single view depth estimation: Geometry to the rescue. In *European Conference on Computer Vision*, pages 740–756. Springer, 2016. 2
- [17] Clément Godard, Oisín Mac Aodha, and Gabriel J Brostow. Unsupervised monocular depth estimation with left-right consistency. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 270–279, 2017. 2
- [18] Wilfried Hartmann, Silvano Galliani, Michal Havlena, Luc Van Gool, and Konrad Schindler. Learned multi-patch similarity. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 1586–1594, 2017. 2
- [19] Paul Henderson and Vittorio Ferrari. Learning single-image 3d reconstruction by generative modelling of shape, pose and shading. *International Journal of Computer Vision*, pages 1–20, 2019. 2
- [20] Po-Han Huang, Kevin Matzen, Johannes Kopf, Narendra Ahuja, and Jia-Bin Huang. Deepmvs: Learning multi-view stereopsis. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 2821–2830, 2018. 1, 2
- [21] Sunghoon Im, Hae-Gon Jeon, Stephen Lin, and In-Su Kweon. Dpsnet: End-to-end deep plane sweep stereo. In *7th International Conference on Learning Representations, ICLR 2019. International Conference on Learning Representations, ICLR*, 2019. 1, 2
- [22] Eldar Insafutdinov and Alexey Dosovitskiy. Unsupervised learning of shape and pose with differentiable point clouds. In *Advances in Neural Information Processing Systems*, pages 2807–2817, 2018. 2
- [23] Mengqi Ji, Juergen Gall, Haitian Zheng, Yebin Liu, and Lu Fang. Surfacenet: An end-to-end 3d neural network for multiview stereopsis. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 2307–2315, 2017. 1, 2
- [24] Angjoo Kanazawa, Michael J Black, David W Jacobs, and Jitendra Malik. End-to-end recovery of human shape and pose. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 7122–7131, 2018. 2
- [25] Angjoo Kanazawa, Shubham Tulsiani, Alexei A Efros, and Jitendra Malik. Learning category-specific mesh reconstruction from image collections. In *Proceedings of the European Conference on Computer Vision (ECCV)*, pages 371–386, 2018. 2
- [26] Sing Bing Kang, Richard Szeliski, and Jinxiang Chai. Handling occlusions in dense multi-view stereo. In *Proceedings of the 2001 IEEE Computer Society Conference on Com-*

- puter Vision and Pattern Recognition. *CVPR 2001*, volume 1, pages I–I. IEEE, 2001. 2
- [27] Abhishek Kar, Christian Häne, and Jitendra Malik. Learning a multi-view stereo machine. In *Advances in neural information processing systems*, pages 365–376, 2017. 1, 2
- [28] Hiroharu Kato, Yoshitaka Ushiku, and Tatsuya Harada. Neural 3d mesh renderer. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 3907–3916, 2018. 2
- [29] Michael Kazhdan and Hugues Hoppe. Screened poisson surface reconstruction. *ACM Transactions on Graphics (ToG)*, 32(3):29, 2013. 5
- [30] Alex Kendall, Hayk Martirosyan, Saumitro Dasgupta, Peter Henry, Ryan Kennedy, Abraham Bachrach, and Adam Bry. End-to-end learning of geometry and context for deep stereo regression. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 66–75, 2017. 2
- [31] Arno Knapitsch, Jaesik Park, Qian-Yi Zhou, and Vladlen Koltun. Tanks and temples: Benchmarking large-scale scene reconstruction. *ACM Transactions on Graphics (ToG)*, 36(4):78, 2017. 6
- [32] Vladimir Kolmogorov and Ramin Zabih. Multi-camera scene reconstruction via graph cuts. In *European conference on computer vision*, pages 82–96. Springer, 2002. 2
- [33] Kiriakos N Kutulakos and Steven M Seitz. A theory of shape by space carving. *International journal of computer vision*, 38(3):199–218, 2000. 2
- [34] Lubor Ladicky, Olivier Saurer, SoHyeon Jeong, Fabio Maninchedda, and Marc Pollefeys. From point clouds to mesh using regression. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 3893–3902, 2017. 2
- [35] Maxime Lhuillier and Long Quan. A quasi-dense approach to surface reconstruction from uncalibrated images. *IEEE transactions on pattern analysis and machine intelligence*, 27(3):418–433, 2005. 2
- [36] Chen-Hsuan Lin, Chen Kong, and Simon Lucey. Learning efficient point cloud generation for dense 3d object reconstruction. In *Thirty-Second AAAI Conference on Artificial Intelligence*, 2018. 2
- [37] Lars Mescheder, Michael Oechsle, Michael Niemeyer, Sebastian Nowozin, and Andreas Geiger. Occupancy networks: Learning 3d reconstruction in function space. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 4460–4470, 2019. 2
- [38] Despoina Paschalidou, Osman Ulusoy, Carolin Schmitt, Luc Van Gool, and Andreas Geiger. Raynet: Learning volumetric 3d reconstruction with ray potentials. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 3897–3906, 2018. 2
- [39] J-P Pons, Renaud Keriven, O Faugeras, and Gerardo Hermosillo. Variational stereovision and 3d scene flow estimation with statistical similarity measures. In *IEEE 9th International Conference on Computer Vision*, page 597. IEEE, 2003. 2
- [40] Stephan R Richter and Stefan Roth. Matryoshka networks: Predicting 3d geometry via nested shape layers. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 1936–1944, 2018. 2
- [41] Gernot Riegler, Ali Osman Ulusoy, Horst Bischof, and Andreas Geiger. Octnetfusion: Learning depth fusion from data. In *2017 International Conference on 3D Vision (3DV)*, pages 57–66. IEEE, 2017. 2
- [42] Olaf Ronneberger, Philipp Fischer, and Thomas Brox. U-net: Convolutional networks for biomedical image segmentation. In *International Conference on Medical image computing and computer-assisted intervention*, pages 234–241. Springer, 2015. 3
- [43] Johannes L Schönberger, Enliang Zheng, Jan-Michael Frahm, and Marc Pollefeys. Pixelwise view selection for unstructured multi-view stereo. In *European Conference on Computer Vision*, pages 501–518. Springer, 2016. 2
- [44] Steven M Seitz, Brian Curless, James Diebel, Daniel Scharstein, and Richard Szeliski. A comparison and evaluation of multi-view stereo reconstruction algorithms. In *2006 IEEE computer society conference on computer vision and pattern recognition (CVPR’06)*, volume 1, pages 519–528. IEEE, 2006. 2
- [45] Daeyun Shin, Zhile Ren, Erik B Sudderth, and Charless C Fowlkes. 3d scene reconstruction with multi-layer depth and epipolar transformers. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 2172–2182, 2019. 2
- [46] Ayan Sinha, Asim Unmesh, Qixing Huang, and Karthik Ramani. Surfnet: Generating 3d shape surfaces using deep residual networks. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 6040–6049, 2017. 2, 5
- [47] Chengzhou Tang and Ping Tan. Ba-net: Dense bundle adjustment network. *arXiv preprint arXiv:1806.04807*, 2018. 1
- [48] Engin Tola, Christoph Strecha, and Pascal Fua. Efficient large-scale multi-view stereo for ultra high-resolution image sets. *Machine Vision and Applications*, 23(5):903–920, 2012. 2, 5
- [49] Shubham Tulsiani, Tinghui Zhou, Alexei A Efros, and Jitendra Malik. Multi-view supervision for single-view reconstruction via differentiable ray consistency. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 2626–2634, 2017. 2
- [50] Ali Osman Ulusoy, Andreas Geiger, and Michael J Black. Towards probabilistic volumetric reconstruction using ray potentials. In *2015 International Conference on 3D Vision*, pages 10–18. IEEE, 2015. 2
- [51] Benjamin Ummenhofer, Huizhong Zhou, Jonas Uhrig, Nikolaus Mayer, Eddy Ilg, Alexey Dosovitskiy, and Thomas Brox. Demon: Depth and motion network for learning monocular stereo. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 5038–5047, 2017. 2
- [52] Jinglu Wang, Bo Sun, and Yan Lu. Mvpnet: Multi-view point regression networks for 3d object reconstruction from a single image. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 33, pages 8949–8956, 2019. 2
- [53] Nanyang Wang, Yinda Zhang, Zhuwen Li, Yanwei Fu, Wei Liu, and Yu-Gang Jiang. Pixel2mesh: Generating 3d mesh

- models from single rgb images. In *Proceedings of the European Conference on Computer Vision (ECCV)*, pages 52–67, 2018. [2](#)
- [54] Jiajun Wu, Yifan Wang, Tianfan Xue, Xingyuan Sun, Bill Freeman, and Josh Tenenbaum. Marrnet: 3d shape reconstruction via 2.5 d sketches. In *Advances in neural information processing systems*, pages 540–550, 2017. [1](#), [2](#)
- [55] Jiajun Wu, Chengkai Zhang, Xiuming Zhang, Zhoutong Zhang, William T Freeman, and Joshua B Tenenbaum. Learning shape priors for single-view 3d completion and reconstruction. In *Proceedings of the European Conference on Computer Vision (ECCV)*, pages 646–662, 2018. [2](#)
- [56] Zexiang Xu, Sai Bi, Kalyan Sunkavalli, Sunil Hadap, Hao Su, and Ravi Ramamoorthi. Deep view synthesis from sparse photometric images. *ACM Transactions on Graphics (TOG)*, 38(4):1–13, 2019. [2](#), [3](#)
- [57] Yao Yao, Shiwei Li, Siyu Zhu, Hanyu Deng, Tian Fang, and Long Quan. Relative camera refinement for accurate dense reconstruction. In *2017 International Conference on 3D Vision (3DV)*, pages 185–194. IEEE, 2017. [2](#)
- [58] Yao Yao, Zixin Luo, Shiwei Li, Tian Fang, and Long Quan. Mvsnet: Depth inference for unstructured multi-view stereo. In *Proceedings of the European Conference on Computer Vision (ECCV)*, pages 767–783, 2018. [1](#), [2](#), [3](#), [4](#), [5](#), [6](#), [8](#)
- [59] Yao Yao, Zixin Luo, Shiwei Li, Tianwei Shen, Tian Fang, and Long Quan. Recurrent mvsnet for high-resolution multi-view stereo depth inference. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 5525–5534, 2019. [2](#), [3](#), [5](#), [6](#), [8](#)
- [60] Xiuming Zhang, Zhoutong Zhang, Chengkai Zhang, Josh Tenenbaum, Bill Freeman, and Jiajun Wu. Learning to reconstruct shapes from unseen classes. In *Advances in Neural Information Processing Systems*, pages 2263–2274, 2018. [2](#)
- [61] Huizhong Zhou, Benjamin Ummenhofer, and Thomas Brox. Deeptam: Deep tracking and mapping. In *Proceedings of the European Conference on Computer Vision (ECCV)*, pages 822–838, 2018. [2](#)
- [62] Tinghui Zhou, Matthew Brown, Noah Snavely, and David G Lowe. Unsupervised learning of depth and ego-motion from video. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 1851–1858, 2017. [2](#)