

Holistic and Contextual Evidential Stereo-LiDAR Fusion for Depth Estimation

Jiayuan Fan, Haixiang Chen, Weide Liu, Xun Xu, Jun Cheng, *Senior Member, IEEE*

Abstract—Stereo-LiDAR fusion is often used for autonomous systems such as self-driving cars as the two modalities are complementary to each other. Existing Stereo-LiDAR fusion methods are mostly at feature level or outcome level, without considering the uncertainty of the depth estimation in each modality. To this end, we propose a holistic and contextual evidential Stereo-LiDAR fusion network (HCENet) for depth estimation, which considers both intra-modality and inter-modality uncertainties from stereo matching and LiDAR point cloud depth completion. We design a dual network structure that consists of a stereo matching branch and a LiDAR depth completion branch with new introduced uncertainty estimation modules for both two branches. Specifically, a multi-scale depth guided feature aggregation module is first developed to enable information propagation at early input stage, and then followed by fusing the predicted depths from two branches based on evidential uncertainties to generate the final output. Extensive experimental results on KITTI depth completion and Virtual KITTI2 datasets achieve RMSE of 599.3 and 2253.1, and show that our method outperforms state-of-the-art SLFNet by 6.52% and 20.7% respectively.

Index Terms—Stereo-LiDAR fusion, Uncertainty Estimation, Depth Completion, Stereo Matching.

I. INTRODUCTION

Nowadays perceiving 3D information of the scenery becomes an essential task in many applications, such as autonomous vehicles [1], obstacle detection and avoidance [2], [3], mobile robotics [4], etc. Ensuring the high-precision and reliability of depth maps becomes critical especially in safety-critical applications. With the rapid development of active and passive sensors, the LiDAR and stereo depth sensors are often used to measure the depth information. As LiDAR often acquires sparse point cloud, depth completion is often conducted to compute dense depth maps [5]–[9]. However, their performance are still limited to the sparsity of LiDAR point clouds and properties of captured objects. Alternatively, the stereo-based depth estimation methods [10]–[13] that estimate the depth maps from a pair of left and right RGB images are studies in recent years. Different from LiDAR that often captures sparse but accurate depth, stereo matching methods compute dense but less accurate depth [14]. The performance of stereo matching methods are mainly restricted

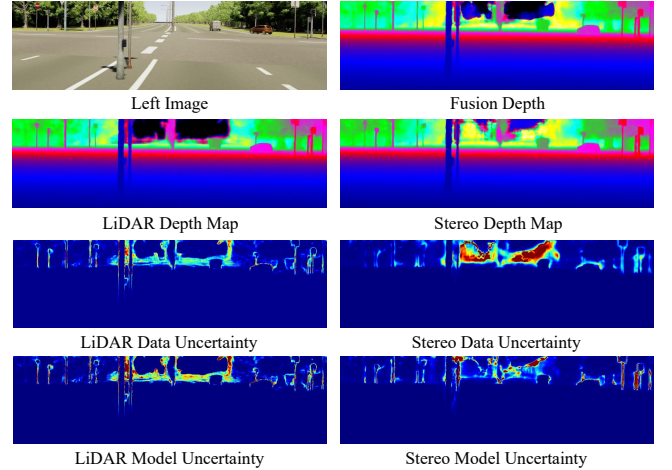


Fig. 1. Our proposed method estimates the data uncertainty and model uncertainty for stereo matching and LiDAR depth completion. The predicted uncertainty maps have large values at the boundaries of the object. We obtain the fusion depth by combining the results from LiDAR depth completion and stereo matching based evidential fusion.

by the error of pixel matching between left and right images, especially for texture-free and repetitive areas. Studies have shown that the LiDAR and stereo matching are complementary to each other. In order to explore the LiDAR and stereo images jointly, many of existing fusion methods [15], [16] integrate the two modalities at the feature level. These Stereo-LiDAR fusion methods achieve some progress on the complementary information exploration. However, the uncertainty of the data has not been fully considered, leading to untrustworthy depth imaging or estimation.

Both the RGB and the point cloud data may suffer from some distortions or corruptions during data acquisition, such as saturation, sensor interference, adversarial attacks, etc. Such distortion or corruption often lead to uncertainty and the untrustworthy depth estimation results. There are two main types of uncertainties including the data uncertainty and model uncertainty. The data uncertainty is also called aleatoric uncertainty and it reflects the degree to which the data is noisy, corrupted and inaccurate, compared with underlying accurate data. The model uncertainty is also named the epistemic uncertainty and it refers to the degree of predicted results of trained models compared to real results, describing the confidence of the model output. As shown in Fig. 1, we estimate data and model uncertainties for both depth completion and stereo matching branches together as well as the depth maps for trustworthy depth fusion. Quantifying the data and model

This research is supported in part by National Natural Science Foundation of China (No. 62101137), Shanghai Natural Science Foundation (No. 23ZR1402900) and the Agency for Science, Technology and Research (A*STAR) under its AME Programmatic Funds (Grant No. M23L7b0021).

J. Fan and H. Chen are with Academy for Engineering and Technology, Fudan University, Shanghai, China (email: jyfan@fudan.edu.cn, 21210860038@m.fudan.edu.cn).

W. Liu, X. Xun and J. Cheng are with Institute for Infocomm Research (I²R), Agency for Science, Technology and Research (A*STAR), Singapore 138632 (email: {liu_weide, xu_xun, cheng_jun}@i2r.a-star.edu.sg).

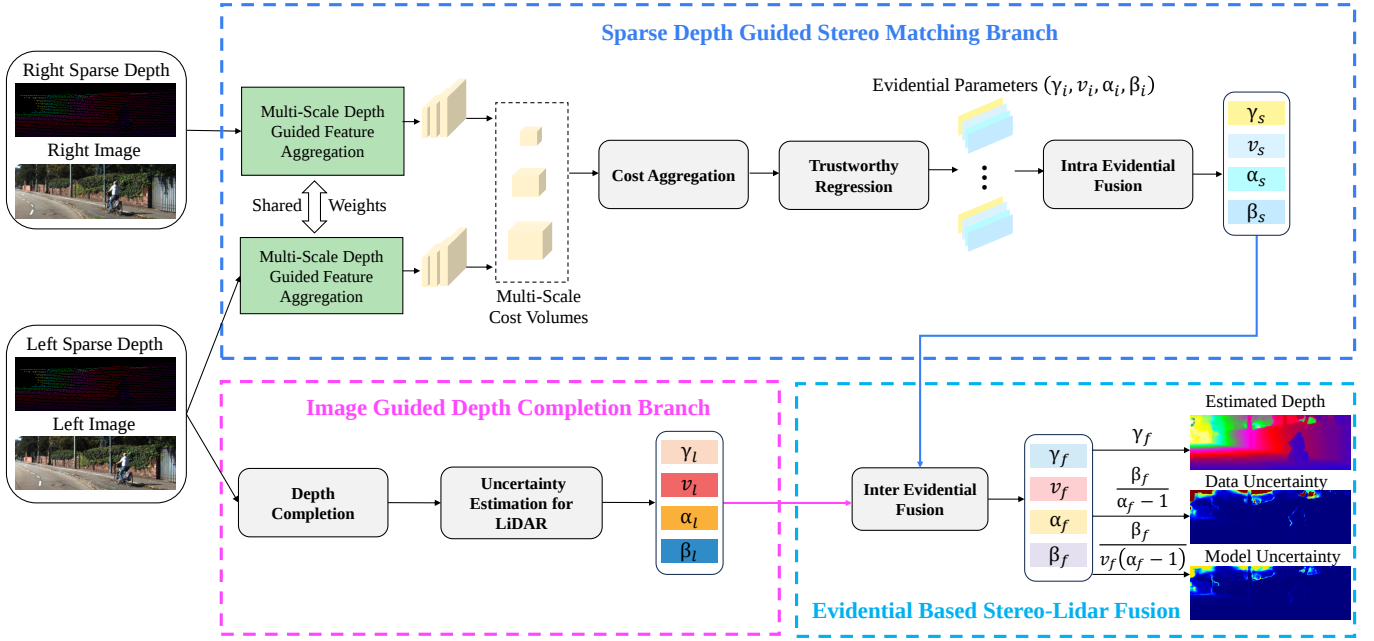


Fig. 2. The overview of our proposed method. It mainly consists of sparse depth guided stereo matching branch and image guided depth completion branch, followed by the evidential based Stereo-LiDAR fusion for final dense depth generation. In the sparse depth guided stereo matching branch, two pairs of LiDAR and stereo images are fed into multi-scale depth guided feature aggregation for estimated dense depth and stereo uncertainty estimation. In the image guided depth completion branch, a pair of LiDAR and stereo images are used for estimated dense depth and LiDAR uncertainty estimation.

uncertainty is practical for enhancing the reliability of depth estimation.

With these motivations in mind, we propose a holistic and contextual evidential Stereo-LiDAR fusion network (HCENet) to kill two birds with one stone (see Fig. 2). It employs a dual network structure including a stereo matching branch and a LiDAR depth completion branch. The framework enables uncertainty estimations in the two branches for fusion and elegantly integrate multi-modal information for evidential based depth fusion. To enable information propagation at early stage, we also propose a multimodal feature aggregation module at the early input stage.

Our main contributions are summarized as follows:

- We propose a novel holistic and contextual evidential based Stereo-LiDAR fusion for depth estimation, which considers the uncertainties and enables both intra-evidential and inter-evidential fusion.
- We propose a feature aggregation module to propagate the information from LiDAR point cloud to camera image for cost volume construction. It exploits complimentary information from two modalities to improve depth estimation.
- Extensive experiments demonstrate that our proposed method outperforms the state-of-the-art Stereo-LiDAR fusion methods on KITTI depth completion and Virtual KITTI2 datasets.

II. RELATED WORK

A. Stereo matching

With the advent and maturation of deep learning, contemporary approaches to stereo matching have predominantly embraced a learning-based foundation. These methods aggregate

the features of the left and right images to construct the cost volume and then calculate the disparity value through 3D convolution [17]–[24]. The disparity can then be converted into depth value through the focal length and baseline of the camera. Some works construct cost volume in different ways. Kendall et al. [25] and Chang et al. [10] construct cost volume using concatenation. Guo et al. [26] simultaneously construct a group-wise correlation volume and a concatenation volume. Xu et al. [27] introduce a sparse points-based intra-scale and cross-scale cost aggregation method to replace 3D convolution, which significantly improves the speed. Most of these methods need to set a maximum disparity according to the datasets. Li et al. [28] exploit cross-attention for matching pixels and does not require a pre-defined disparity range. But it requires additional occlusion masks for training, which is not available for most dataset. Shen et al. [29] introduce a volume fusion module to directly combine multi-scale 4D volumes and calculate a multi-level loss to accelerate the convergence of the model. In this paper, we adopt PCWNet [29] as a baseline to combine with depth from LiDAR.

B. Depth Completion

LiDAR can provide accurate but sparse depth. To generate dense depth map from a sparse depth map, depth completion is often conducted. Many methods [30]–[35] have been proposed for LiDAR depth completion with the guidance of a reference color image to integrate structural information of objects. Typically, early fusion methods concatenate features from RGB images and sparse depth for depth completion. Yang et al. [36] introduce a late fusion strategy which infers the posterior distribution of a dense depth map associated with an image. Forkel et al. [37] introduce a real-time fusion method

that integrates LiDAR point clouds into stereo processing using the semi-global matching algorithm. Hu et al. [38] introduce a dual-branch backbone to generate color- and depth-dominant depth maps, then fuse them to get the final output. Zhao et al. [39] introduce a mechanism of neighboring attention. Rho et al. [40] introduce transformer architecture for depth completion and introduced a guided-attention block to fuse depth and color features. Zhang et al. [41] integrate convolutional layers and transformers into a single block for deep completion, enabling the encoding of both local and global information simultaneously. To avoid large amount of parameters in transformer based approaches, we chose to use PENet [38] in the depth completion branch. From the models, we further introduced uncertainty prediction on top of PENet. Our method, as compared to existing depth completion approaches, incorporates uncertainty estimation for the completed depth map. The resulting depth map, along with its associated uncertainty, can be fused with the depth map from the stereo modality.

C. Uncertainty Estimation

The data uncertainty can be caused by the noise of training data, the model uncertainty represents the shortcomings of the network. Ensemble [42] methods predict model uncertainty by training multiple models with different initializations and modeling the distribution over parameters. Monte Carlo (MC) Dropout [43] obtains multiple different prediction distributions by randomly turning off neurons in the neural network. Both Ensemble and MC Dropout require multiple forward passes, depth values and uncertainties can be obtained by computing the mean and variance of multiple predicted values respectively.

Recently, trustworthy machine learning has been used to compute uncertainty for evidential based multi-view classification [44], which shows significant improvement. Many works use uncertainty for depth prediction. Gansbeke et al. [45] integrate global and local information to generate the depth map via confidence weights. Cheng et al. [12] calculate adaptive thin volume with an uncertainty-aware cascaded design. Shen et al. [46] adaptively adjust the disparity search range. In contrast to the previous methods that rely on variance-based uncertainty, deep evidential learning [47] is employed to estimate the uncertainties of the disparity map [48]. Lou et al. [49] utilize evidential based approach to fuse the outputs from different stereo matching methods.

Based on evidential learning [47], every depth estimation d is drawn from a normal distribution with unknown mean and variance (μ, σ^2) , which are assumed to be drawn from normal and inverse-gamma distributions respectively:

$$d \sim \mathcal{N}(\mu, \sigma^2), \mu \sim \mathcal{N}(\gamma, \sigma^2 v^{-1}), \sigma^2 \sim \Gamma^{-1}(\alpha, \beta), \quad (1)$$

where $\Gamma(\cdot)$ denotes the gamma function, $\gamma \in \mathbb{R}$, $v > 0$, $\alpha > 1$, $\beta > 0$.

With the assumption that the mean and the variance are independent, the posterior distribution is formulated as a normal-inverse gamma distribution $\text{NIG}(\gamma, v, \alpha, \beta)$. In this work, we follow the evidential learning approach and predict

uncertainties for both stereo matching branch and depth completion branch for fusion. Different from [48], we integrate the uncertainty computation with multi-scale cost volume for an intra-modality fusion. Moreover, we also incorporate inter-modality fusion between stereo matching and LiDAR based depth completion.

D. Stereo-LiDAR Fusion

There have been some research endeavors focusing on the fusion of stereo and LiDAR modalities, aiming to exploit their complementary properties. In the method introduced by Zhang et al. [50], a straightforward concatenation of features derived from stereo images and point clouds was utilized for the purpose of depth estimation. Maddern et al. use a probabilistic approach to real-time LiDAR stereo fusion for low-cost 3D perception [15]. However, this probabilistic fusion approach performs poorly in scenarios with limited depth information. To address this issue, Park et al. [16] employ a CNN model to fuse the input Lidar disparity and Stereo disparity. Afterwards, LiDAR and stereo disparities are utilized to generate calibration parameters [51], which are updated every frame for the purpose of registering LiDAR and stereo disparities. Wang et al. [52] fuse LiDAR information into a stereo matching network at an early stage. Choe et al. [53] extract image and point cloud features and fuse two modalities in a depth cost volume to encode semantic and geometric information. Cheng et al. [54] tackle the noise and misalignment in the data by a feedback loop. Xu et al. [55] introduce a novel lightweight scheme that combines RGB information and sparse LiDAR points to generate a semi-dense depth map and employs a varying-weight Gaussian guiding method for effective cost volume aggregation. He et al. [56] introduce a novel framework that integrates semantic information from stereo images and spatial information from raw point clouds for enhanced 3D object detection by utilizing a residual attention learning mechanism. Our proposed method stands out from existing Stereo-LiDAR fusion techniques by not only integrating different modalities at the feature level but also leveraging uncertainty to fuse multiple depth maps.

III. METHODOLOGY

The overview of our proposed method is shown in Fig. 2. It consists of two main branches. Specifically, one branch is designed for sparse depth guided stereo matching and the other one is designed for image guided depth completion, where each branch predicts the depth map as well as modality-specific uncertainties. Finally, the outputs from the two branches are combined via evidential based Stereo-LiDAR fusion.

In the sparse depth guided stereo matching branch, we first apply a multi-scale depth guided feature aggregation on the stereo images (i.e., the left and right images) and their corresponding LiDAR depth maps. Subsequently, a stereo matching network is employed, whereby the adoption of the architecture proposed by the recent PCWNet [29] is predicated upon its commendable efficacy in facilitating cross-domain generalization as well as enhancing stereo matching accuracy. By

performing multi-scale cost volume and disparity regression, the estimated dense depth can be obtained by converting the disparity into depth. We integrate the uncertainty computation into the architecture to compute the uncertainty of the stereo matching results. In the depth completion branch, the left image and corresponding left sparse depth are used for deep completion network to generate the estimated dense depth. We adopt the network architecture from [38]. Similar to that in stereo matching, uncertainty of depth completion is also computed. After the depth map and uncertainty map of the two branches are obtained, we can take the pixel-wise uncertainty into account and fuse the evidential distribution of pixels of two different modalities for final depth. In this way, we can also get the model uncertainty and data uncertainty of the fused depth map.

A. Sparse Depth Guided Stereo Matching

In stereo matching, the process typically begins with feature extraction and cost volume construction. The subsequent steps in stereo matching, namely 3D cost aggregation and disparity regression, utilize the information in the cost volume to produce the final disparity map. Cost volume encodes important features for subsequent disparity calculations. Different from only using stereo images to construct cost volumes [29], we here enrich the cost volumes by combining the information of the from both sparse depth and stereo images using a new proposed multi-scale depth guided feature aggregation module.

1) *Multi-scale Depth Guided Feature Aggregation*: Inspired by multi-modality features fusion method [57], we employ a multi-scale depth guided feature aggregation strategy to incorporate depth information into RGB image. As shown in Fig. 3(a). Given an image I and corresponding sparse depth D , initially, we employ separate convolutional layers to generate initial features for both image I and depth map D , which serve as inputs to the initial multimodal feature aggregation (MFA) module. Subsequently, the output of each MFA layer is propagated as input to the subsequent MFA layer. A five-stage MFA extraction module is proposed to compute multi-scale feature maps, which is used for cost volume generation subsequently.

As shown in Fig. 3(b), initially, we fuse RGB features f_I and sparse depth features f_D by concatenating them, thereby blending their distinctive attributes at a specific spatial location. Afterward, we use different convolutional layers to generate spatial-wise gate for f_I and f_D , respectively. After obtaining these two gates $G_I \in \mathbb{R}^{1 \times H \times W}$, and $G_D \in \mathbb{R}^{1 \times H \times W}$, a softmax function is used to generate weights for the fusion between RGB features and sparse depth features as follows:

$$A_I = \frac{e^{G_I}}{e^{G_I} + e^{G_D}}, \quad (2)$$

$$A_D = \frac{e^{G_D}}{e^{G_I} + e^{G_D}}. \quad (3)$$

An intermediate feature f_w is then computed by weighting the RGB and depth features as follows:

$$f_w = f_I \cdot A_I + f_D \cdot A_D. \quad (4)$$

Then we compute f'_I and f'_D as mean of f_w and RGB features or depth features respectively, which are used for next stage of MFA:

$$f'_I = \frac{f_w + f_I}{2}, \quad (5)$$

$$f'_D = \frac{f_w + f_D}{2}. \quad (6)$$

We concatenate f_w and f_I to get f_A :

$$f_A = [f_w, f_I]. \quad (7)$$

After the multi-scale depth-guided feature aggregation, three feature maps with sizes of $\frac{H}{4} \times \frac{W}{4}$, $\frac{H}{8} \times \frac{W}{8}$, and $\frac{H}{16} \times \frac{W}{16}$ are obtained for constructing the multi-scale cost volume. H and W are the height and width of the initial image.

2) *Uncertainty Estimation in Stereo Matching*: We adopt the recent pyramid combination and warping network (PCWNet) [29] because of its good performance on both cross-domain generalization and stereo matching accuracy. PCWNet constructs combination volumes on the upper levels of the pyramid and develop a cost volume fusion module to integrate them for initial disparity estimation. Then, a warping volume is constructed on the last level of the pyramid and employed to refine the initial disparity. Slightly different from the original PCWNet, we use three scales instead of four to reduce computational cost for later fusion with LiDAR, and we do not employ depth refinement. By performing multi-scale cost volume and disparity regression, the estimated dense depth is obtained. In this work, we compute the uncertainties for the multi-scale cost volume to achieve intra-scale fusion.

To estimate the parameters of NIG distribution rather than merely predicting the disparity map, we modify the disparity regression module into a trustworthy regression with multi-channel output while keeping the remaining modules unchanged. As shown in Fig. 4, the proposed trustworthy regression module uses a single branch of 3D convolution and the up-sampling block to compute a 4-channel output $O \in \mathbb{R}^{D \times H \times W \times 4}$, where D is the maximum of disparity, H and W are the height and width of the stereo images. We set $D = 192$ in our implementation. The distribution parameters can be computed as follows:

$$O_\gamma, O_v, O_\alpha, O_\beta = \text{Split}(O, \text{dim} = -1), \quad (8)$$

$$p = \text{Softmax}(O_\gamma), \quad (9)$$

$$\gamma = \sum_{k=0}^D k \cdot p_k, \quad \ell_i = \sum_{k=0}^D O_i \cdot p_k, \quad (10)$$

where k denotes disparity level, p denotes the probability, and $i \in \{v, \alpha, \beta\}$. Finally, a Softplus activation is applied on the ℓ_i to obtain v, α, β .

After obtaining the evidence distribution parameters $(\gamma, v, \alpha, \beta)$ of each pixel, predicted depth value, data and model uncertainties can be calculated as follow:

$$\underbrace{\mathbb{E}[\mu]}_{\text{prediction}} = \gamma, \quad (11)$$

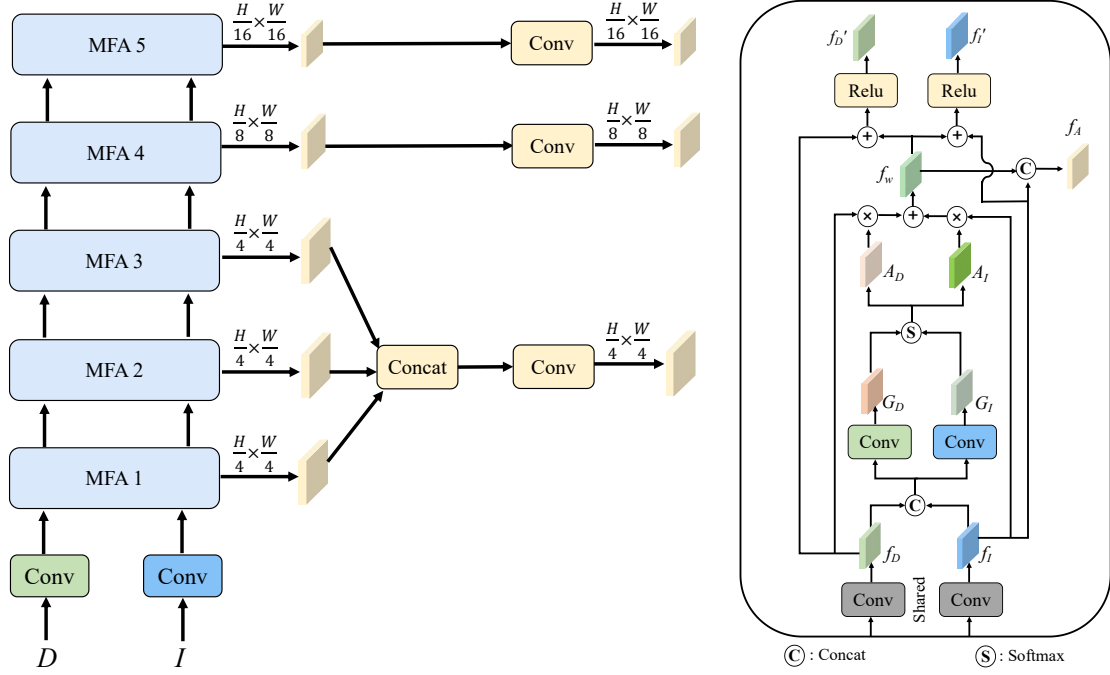


Fig. 3. Schematic diagram of the five-stage multi-scale depth guided feature aggregation. (a) Diagram of the multi-scale depth guided feature aggregation (b) Diagram of the multimodal feature aggregation (MFA) module.

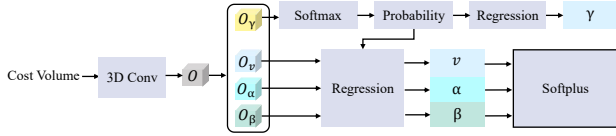


Fig. 4. The illustrates the trustworthy regression module designed for stereo matching.

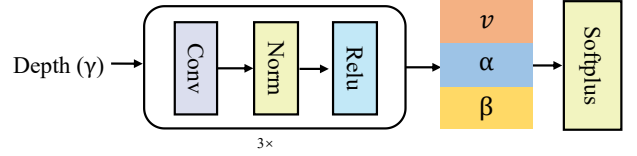


Fig. 5. The network structure for uncertainty estimation in depth completion.

$$\underbrace{\mathbb{E}[\sigma^2]}_{\text{data uncertainty}} = \frac{\beta}{\alpha - 1}, \quad (12)$$

$$\underbrace{\text{Var}[\mu]}_{\text{model uncertainty}} = \frac{\beta}{v(\alpha - 1)}. \quad (13)$$

Parameters v affects model uncertainty while β and α affect both data uncertainty and model uncertainty.

B. Uncertainty Aware Image Guided Depth Completion

We follow the main network structure in the precise and efficient image guided depth completion network (PENet) [38] and modify it for uncertain aware image guided depth completion. PENet uses a two-branch backbone to exploit color-dominant and LiDAR-dominant information for two indepent depth maps. The two dense depth maps are then fused using the same strategy as in FusionNet [45]. We utilize the PENet while excluding the utilization of its depth refinement module.

Fig. 5 illustrates network structure of the uncertainty computation. On top of the structure from PENet, we denote its estimated depth value as γ . We further use three convolutional layers with the kernel size 3×3 followed by BatchNorm and

Relu to generate per-pixel uncertainty parameters v , α and β , which enable us to compute the data and model uncertainties in depth completion branch.

C. Fusion Strategy based on Evidence

Simply taking a weighted average of the depths from different scales or modalities may result in mutual interference among the different depth maps. Therefore, it is essential to incorporate the uncertainty of each depth map for fusion. We adopt the fusion strategy with the mixture of NIG distribution (MoNIG) [58] for its excellent mathematical properties to perform both intra evidential fusion and inter evidential fusion. Specifically, given k sets of parameters of NIG distributions, the MoNIG distribution can be computed with the following operations:

$$\text{MoNIG}(\gamma, v, \alpha, \beta) = \text{NIG}(\gamma_1, v_1, \alpha_1, \beta_1) \oplus \text{NIG}(\gamma_2, v_2, \alpha_2, \beta_2) \oplus \dots \oplus \text{NIG}(\gamma_k, v_k, \alpha_k, \beta_k), \quad (14)$$

where $\oplus$ represents the summation operation of two NIG distributions:

$$\text{NIG}(\gamma, v, \alpha, \beta) \triangleq \text{NIG}(\gamma_m, v_m, \alpha_m, \beta_m) \oplus \text{NIG}(\gamma_n, v_n, \alpha_n, \beta_n), \quad (15)$$

where $\gamma = \frac{v_m \gamma_m + v_n \gamma_n}{v_m + v_n}$, $\alpha = \alpha_m + \alpha_n + \frac{1}{2}$, $v = v_m + v_n$, $\beta = \beta_m + \beta_n + \frac{1}{2} v_m (\gamma_m - \gamma)^2 + \frac{1}{2} v_n (\gamma_n - \gamma)^2$.

1) *Intra Evidential Fusion of Multi-scale Cost Volume Stereo Matching*: Multi-scale cost volume is often used in stereo matching to exploit the features from different layers of the feature extractor. We construct three levels of the cost volumes with $\frac{1}{16}$, $\frac{1}{8}$, $\frac{1}{4}$ scaled features and employ the cost volume fusion module [29] to early combine the knowledge from multi-scale receptive field. These feature maps contain coarse-to-fine semantic information such as textures, boundaries, and regions. Then we apply three branches of trustworthy regression modules to generate parameters of NIG distributions $(\gamma_i, v_i, \alpha_i, \beta_i)$, where $i \in \{1, 2, 3\}$. Intra evidential fusion module integrates three NIG distributions into one NIG distribution $NIG(\gamma_s, v_s, \alpha_s, \beta_s)$:

$$NIG(\gamma_s, v_s, \alpha_s, \beta_s) = NIG(\gamma_1, v_1, \alpha_1, \beta_1) \oplus \dots \oplus NIG(\gamma_3, v_3, \alpha_3, \beta_3). \quad (16)$$

The intra evidential fusion strategy integrates multi-scale features for trustworthy stereo matching.

2) *Inter Evidential Fusion of LiDAR Depth Completion and Stereo Matching*: Besides the intra evidential fusion of stereo matching, we further apply inter evidential fusion from LiDAR and stereo matching. Denote the NIG distribution of depth from LiDAR as $NIG(\gamma_l, v_l, \alpha_l, \beta_l)$. We fuse it with the depth from stereo matching as follows to get:

$$NIG(\gamma_f, v_f, \alpha_f, \beta_f) = NIG(\gamma_s, v_s, \alpha_s, \beta_s) \oplus NIG(\gamma_l, v_l, \alpha_l, \beta_l), \quad (17)$$

where $\gamma_f = (v_s + v_l)^{-1} (v_s \gamma_s + v_l \gamma_l)$, $\alpha_f = \alpha_s + \alpha_l + \frac{1}{2}$, $v_f = v_s + v_l$, $\beta_f = \beta_s + \beta_l + \frac{1}{2} v_s (\gamma_s - \gamma_f)^2 + \frac{1}{2} v_l (\gamma_l - \gamma_f)^2$.

The parameter v_s and v_l represent the confidence of the depth predicted by the stereo and Lidar modalities, respectively. In this way, pixels with high uncertainty have lower weights in the fusion process, and the depth map after fusion is trustworthy and interpretable. The parameters after fusion $(\gamma_f, v_f, \alpha_f, \beta_f)$ also obey the NIG distribution. Therefore data uncertainty and model uncertainty can be calculated directly. The uncertainty map of the fused depth map is obtained by mixing the NIG distribution of the corresponding pixels of the two branches. It does not need parameters for learning and does not require a convolutional layer.

D. Loss Function

1) *Depth Estimation Loss*: We compute the RMSE losses $\mathcal{L}_D$ for depth obtained through stereo matching, depth completion, and the final fused depth, denoted as $\mathcal{L}_s$, $\mathcal{L}_l$ and $\mathcal{L}_f$ respectively. The depth map of stereo branch is obtained by converting the disparity maps using the camera focal length and baseline. The RMSE loss $\mathcal{L}_D$ is calculated based on the ground truth depth map:

$$\mathcal{L}_D = \sqrt{\frac{1}{|V|} \sum_{v \in V} (\hat{D}_v - D_v^{gt})^2}, \quad (18)$$

where V represents the set of pixels with valid depth in the ground truth depth map, and $|V|$ represent the size of the set

V . $\hat{D}_v$ refers to the predicted depth, while D_v^{gt} represents the ground truth depth.

For stereo matching branch, we also compute smooth L_1 loss on three disparity maps, which can be expressed as follows:

$$\mathcal{L}_{disp} = \sum_{j=0}^2 w_j \cdot \mathcal{L}_{smooth L_1}(\hat{d}_j, d_{gt}), \quad (19)$$

where w_j is the weight of the j^{th} prediction of disparity map, we set $w_0 = 0.7$, $w_1 = 0.7$, $w_2 = 1$, empirically. $\hat{d}_j$ and d_{gt} are the j^{th} prediction and ground truth disparity respectively, we only consider pixels with a disparity value greater than zero.

2) *Uncertainty Estimation Loss*: Following [47], we compute a negative logarithm of model evidence $\mathcal{L}^N(\gamma, v, \alpha, \beta)$ and an evidence regularizer $\mathcal{L}^R(\gamma, v, \alpha)$. The model evidence $\mathcal{L}^N(\gamma, v, \alpha, \beta)$ is computed as:

$$\begin{aligned} \mathcal{L}^N(\gamma, v, \alpha, \beta) &= \frac{1}{2} \log\left(\frac{\pi}{v}\right) - \alpha \log(\Omega) \\ &+ \left(\alpha + \frac{1}{2}\right) \log((D^{gt} - \gamma)^2 v + \Omega) \\ &+ \log\left(\frac{\Gamma(\alpha)}{\Gamma(\alpha + \frac{1}{2})}\right), \end{aligned} \quad (20)$$

where $\Omega = 2\beta(1 + v)$.

The evidence regularizer $\mathcal{L}^R(\gamma, v, \alpha)$ is computed as:

$$\mathcal{L}^R(\gamma, v, \alpha) = |D^{gt} - \gamma| \cdot (2v + \alpha). \quad (21)$$

The evidence regularizer encourages large uncertainty. When the difference between the predicted value γ and the ground truth D^{gt} is large. This can be easily seen as the minimization of the term in Eq. (21) would lead to smaller values of v and α and therefore large uncertainties in Eq. (12) and Eq. (13) subsequently.

The total loss $\mathcal{L}^{unc}$ comprises two loss terms that serve the purpose of maximizing and regularizing evidence:

$$\mathcal{L}^{unc}(\gamma, v, \alpha, \beta) = \mathcal{L}^N(\gamma, v, \alpha, \beta) + \lambda \mathcal{L}^R(\gamma, v, \alpha), \quad (22)$$

where $\lambda > 0$ controls the balance between the two items. We set $\lambda = 0.01$ in our experiments.

We apply the equations above on the depth maps generated by stereo matching branch, the LiDAR depth completion branch, and the fused depth map to get their corresponding uncertainty losses $\mathcal{L}^{unc}(\gamma_s, v_s, \alpha_s, \beta_s)$, $\mathcal{L}^{unc}(\gamma_l, v_l, \alpha_l, \beta_l)$ and $\mathcal{L}^{unc}(\gamma_f, v_f, \alpha_f, \beta_f)$ respectively.

To optimize the disparity value, predicted depth value and the corresponding uncertainty value, the overall loss function is computed as follows:

$$\begin{aligned} \mathcal{L}_{Total} &= \mathcal{L}_l + \mathcal{L}_s + \mathcal{L}_f + \mathcal{L}_{disp} + \mathcal{L}^{unc}(\gamma_s, v_s, \alpha_s, \beta_s) \\ &+ \mathcal{L}^{unc}(\gamma_l, v_l, \alpha_l, \beta_l) + \mathcal{L}^{unc}(\gamma_f, v_f, \alpha_f, \beta_f). \end{aligned} \quad (23)$$

IV. EXPERIMENTS

A. Datasets and Evaluation Metrics

We have implemented the proposed network and tested its performance with the KITTI depth completion dataset [5] and the Virtual KITTI2 dataset [59].

TABLE I

COMPARISON WITH OTHER METHODS ON THE KITTI DEPTH COMPLETION VALIDATION SET. S REFERS TO STEREO IMAGES, L REFERS TO LiDAR AND M REFERS TO MONOCULAR IMAGES. ↓ INDICATES THAT THE SMALLER VALUE THE BETTER PERFORMANCE.

Methods	Modality	RMSE (mm)↓	MAE (mm)↓	IRMSE (1/km)↓	IMAE (1/km)↓
ACMNet [39]	M+L	1037.3	236.4	2.38	0.91
NLSPN [62]	M+L	868.8	236.7	2.61	1.02
GuideNet [63]	M+L	857.8	234.0	2.36	0.99
PENet [38]	M+L	761.3	210.9	2.19	0.90
LiStereo [50]	S+L	832.1	283.9	2.19	1.10
CCVN [52]	S+L	749.3	252.5	1.39	0.80
VPN [53]	S+L	636.2	205.1	1.87	0.98
SLFNet [61]	S+L	641.1	197.0	1.77	0.87
Ours	S+L	599.3	190.0	1.43	0.78

The KITTI depth completion dataset is a large-scale dataset capturing real-world driving scenarios in outdoor environments. It comprises 138 scenes, encompassing a training set of 42,949 images and a validation set of 3,426 images. Because there is no ground truth in the test set, we follow CCVN [52] and use the same 1000 pictures in the validation set for comparison with other methods. The initial resolution is 1242×352 . Due to the absence of ground truth in the top portion of the initial images, a cropping approach is employed during the training process, resulting in images with dimensions of 1216×256 . The depth maps provided by Velodyne HDL-64e are sparse, with ground truth depth generated by combining consecutive 11 frames to form a single depth map, resulting in approximately 30% of pixels containing depth information [52].

Virtual KITTI2 dataset is a synthetic dataset that provides ground truth depth for all pixels and is a more realistic version of the original Virtual KITTI 1.3.1 [60]. Following [61], we use “Scene01” and “Scene02” for training, “Scene06”, “Scene18” and “Scene20” for evaluation. For each scene, we take the only scenario (15-deg-left) for training and evaluation. The training set has 680 images, and the test set has 1446 images.

Following [52], we compute root mean square error (RMSE), mean absolute error (MAE), inverse root mean square error (IRMSE), and inverse mean absolute error (IMAE) to evaluate the performance of our method.

B. Implementation Details

Our method is implemented with PyTorch. The training phase uses the Adam ($\beta_1 = 0.9, \beta_2 = 0.999$) optimizer. We use four NVIDIA GTX 3090 GPUs for experiments, and set the batch size to 4. The initial learning rate is set to $1e-3$. after training for 10 epochs, it decreases to $5e-4$. The training epoch is set to 15 for KITTI depth completion dataset. We use the weights pre-trained on the KITTI depth completion dataset to fine-tune on Virtual KITTI2 dataset for 14 epochs. The maximum depth is set to 100 and 80 on KITTI depth completion dataset and Virtual KITTI2 dataset, respectively.

TABLE II

COMPARISON WITH STEREO-LiDAR FUSION METHODS ON THE VIRTUAL KITTI2 DATASET.

Methods	RMSE(mm)↓	MAE(mm)↓
CCVN [52]	3726.8	915.6
VPN [53]	3217.1	712.0
SLFNet [61]	2843.1	696.2
Ours	2253.1	670.1

TABLE III

PERFORMANCE OF DIFFERENT FEATURE FUSION STRATEGIES ON KITTI DEPTH COMPLETION VALIDATION SET.

Methods	RMSE↓	MAE↓	IRMSE↓	IMAE↓
Concat	614.4	193.4	1.49	0.81
Ours	599.3	190.0	1.43	0.78

C. Comparison With The State-of-the-art Methods

1) *Performance Comparison in KITTI*: We compare our method with four stereo-LiDAR fusion methods including LiStereo [50], CCVN [52], VPN [53], SLFNet [61] in KITTI depth completion dataset. We obtain the results from [61]. Table I summaries the results. Our method is compared with other methods in four metrics. Specifically, it outperforms VPN and SLFNet by 5.8% and 6.5% in RMSE, 7.3% and 3.5% in MAE, 23.5% and 19.2% in IRMSE, 20.4% and 10.3% in IMAE, respectively.

Fig. 6 shows a visual comparison of the results of four samples by the proposed method with those by CCVN and SLFNet. As we can see from the comparison, our method predicts the depth more accurately in regions with objects such as railing, traffic sign board etc. In contrast, CCVN and SLFNet produce inferior predictions where the shapes of the objects are less well preserved from the backgrounds in depth map.

Besides the Stereo-LiDAR fusion methods, we also show the results from monocular depth completion fusion (M+L) methods (ACMNet [39], NLSPN [62], GuideNet [63], PENet [38]) in Table I for comparison. Compared to monocular depth completion methods, our approach also surpasses all the methods with large margins.

2) *Performances in other dataset*: Although KITTI depth completion dataset is a real-world dataset, it may have some limitations as the ground truth of the depth is derived by aggregating LiDAR points from surrounding frames. Therefore, it may suffer from errors for regions such as the glass on cars, occluded areas, etc. Therefore, we also conduct the experiments on Virtual KITTI2 dataset. Since the Virtual KITTI2 dataset only provides stereo image pairs and dense depth pairs, we use randomly generated masks and integrate them with ground truth to obtain sparse depth maps, each with about 5k pixels. Table II shows the comparison on the Virtual KITTI2 dataset. Our method achieves the best performance as well. It outperforms CCVN, VPN and SLFNet in RMSE by 39.5%, 30.0%, 20.7%, respectively.

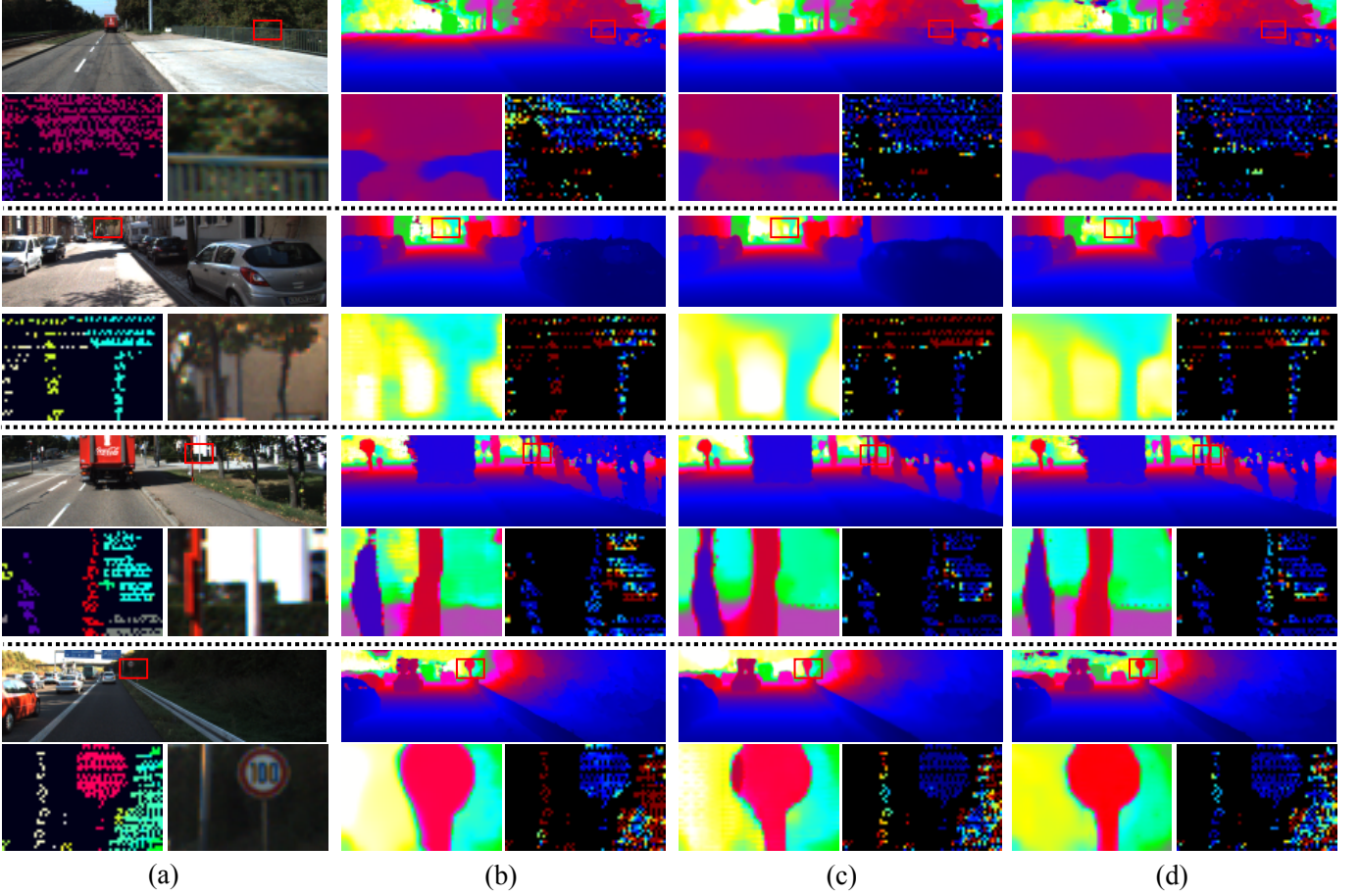


Fig. 6. Visual comparison of results (a) The left image, along with the ground truth depth map and the zoomed-in region (rectangular region in red) for both the ground truth and the image. (b) The results by CCVN. (c) The results by SLFNet. (d) The results by our proposed HCENet.

D. Ablation Study

In this section, we perform ablation studies to assess the impact of different design choices in our network. We first evaluate the benefits of the MFA module. We compare MFA with simple concat. Table III shows the performance of different feature fusion strategies on KITTI depth completion dataset. Our MFA module outperforms simple concat and proves that our fusion method is suitable for feature fusion of sparse depth and RGB image.

Next, we evaluate how different λ affect the performance, which is used to control the balance between model evidence and evidence regularizer in Eq (22). Table IV shows the results for different λ . As we can see, $\lambda = 0.01$ leads to best results. The results show that when there is no penalty ($\lambda = 0$), the performance of the model will be worse, which justifies the needs of the evidence regularization. On the other side, large values would not help much.

We continue the ablation study to justify the effectiveness of intra evidential fusion, the MFA and the MoNIG. Table V shows the results. As can be seen, if the MFA module is not applicable, the RMSE will increase, which shows that the introduction of sparse depth features in the feature fusion stage can help our stereo matching branch to generate more accurate depths. This is because there are many textureless

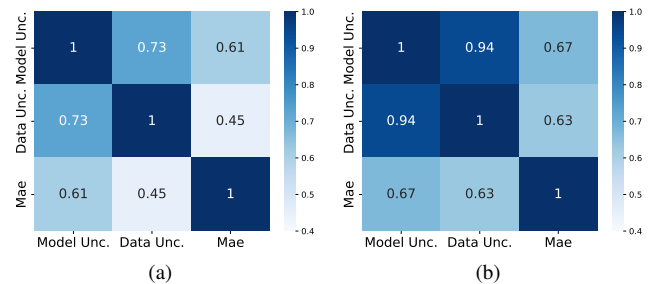


Fig. 7. The Pearson correlation coefficients between the uncertainties and the Mae of estimated depth. (a) The KITTI depth completion dataset. (b) The Virtual KITTI2 dataset.

TABLE IV
COMPARISON OF THE STRENGTH OF EVIDENCE REGULARIZATION ON KITTI DEPTH COMPLETION DATASET.

λ	RMSE↓	MAE↓	IRMSE↓	IMAE↓
0	602.6	208.9	1.43	0.82
0.5	644.0	198.3	1.49	0.82
0.02	600.6	192.6	1.44	0.79
0.01	599.3	190.0	1.43	0.78

and occluded regions in the dataset, and the introduction of LiDAR information can reduce false predictions. Intra evidential fusion and MoNIG both contribute to improvement of the

TABLE V
ABLATION EXPERIMENTS WITH DIFFERENT MODULE COMBINATIONS ON KITTI DEPTH COMPLETION VALIDATION DATASET.

Intra Evidential Fusion	MFA	MoNIG	RMSE↓	MAE↓	IRMSE↓	IMAE↓
✓	✓		651.1	224.1	1.78	0.97
✓		✓	621.6	219.4	1.39	0.82
	✓	✓	604.4	203.8	1.49	0.84
✓	✓	✓	599.3	190.0	1.43	0.78

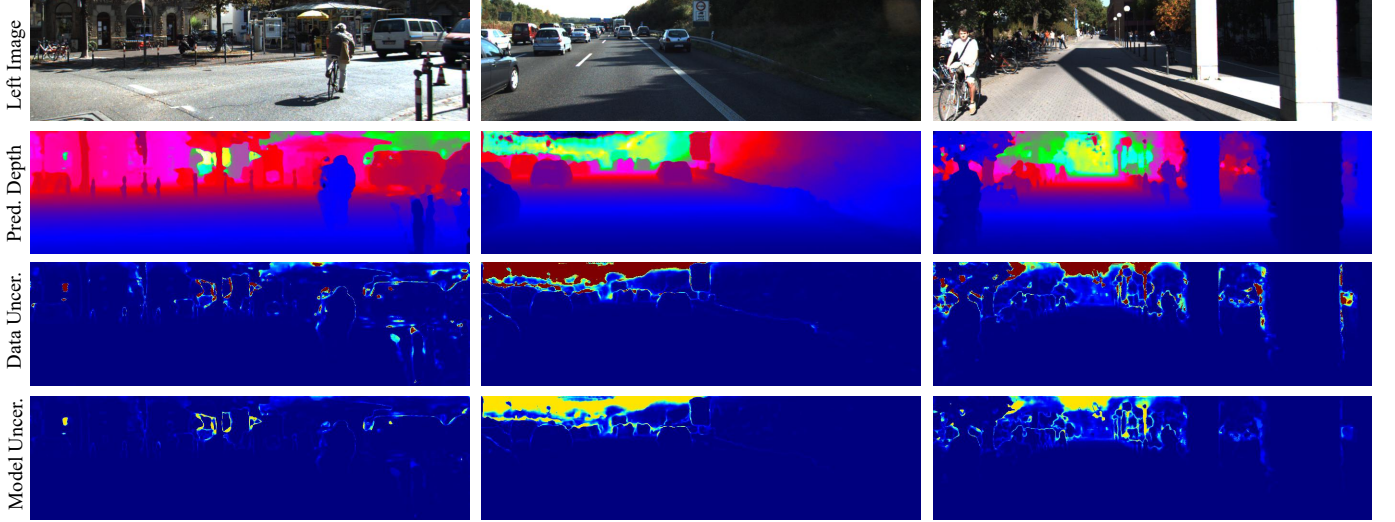


Fig. 8. The top row in the illustration represents the left image, while the second row depicts the fused depth map. The third and fourth rows correspond to the data uncertainty and model uncertainty, respectively. Both data uncertainty and model uncertainty exhibit higher values in the boundary regions of objects and at farther distances.

TABLE VI
RESULTS OF DIFFERENT DEPTH FUSION STRATEGIES ON KITTI DEPTH COMPLETION VALIDATION DATASET.

Methods	RMSE↓	MAE↓	IRMSE↓	IMAE↓
Average	651.1	224.1	1.78	0.97
Relative	625.6	223.0	1.43	0.84
Ours	599.3	190.0	1.43	0.78

model, resulting in a respective decrease of 0.8% and 7.9% in the RMSE. Furthermore, when these three modules are used in conjunction, the model exhibits its best performance across all four metrics.

E. Comparison of Depth Map Fusion Strategies

We also conduct experiments to evaluate the performance enhancements introduced by our depth fusion strategies. In Table VI, we compare different ways to fuse the depth values predicted by the two branches of stereo matching and depth completion. Using the average depth of two branches lacks uncertainty incorporation. It simply combines the depth maps without considering uncertainty and divides the result by two, leading to poor performance. We also compute the relative model uncertainties obtained from two branches by concatenating their model uncertainties and applying the softmax function. This process allows us to obtain the relative uncertainties, which is then used for depth map fusion. Our fusion method is better than these two compared strategies,

which shows that our predictive level fusion can fully exploit the uncertainty of two modalities.

F. Uncertainty Analysis

The deep evidential learning of uncertainty computes both data and model uncertainties. However, data uncertainty cannot be reduced through model improvement, whereas model uncertainty reflects the predictive capability of the model. We computed the Pearson correlation coefficients among the MAE metric and the data and model uncertainties tested on the KITTI depth completion dataset and the Virtual KITTI2 dataset in Fig. 7. The correlation coefficients between model uncertainty and MAE are high on both datasets, with values of 0.61 and 0.67, respectively. The positive correlation coefficient indicates that our predicted uncertainty is able to effectively reflect the errors. Therefore, during inference, the accuracy of predictions can be inferred based on their associated uncertainty. On the two datasets, the correlation coefficients between model uncertainty and data uncertainty are 0.73 and 0.94, respectively. The model uncertainty and data uncertainty are highly correlated, indicating that data uncertainty plays a major role to affect the model uncertainty. We also provide a visualization of the data uncertainty and model uncertainty on KITTI depth completion dataset in Fig. 8 from three samples. As observed, the data and model uncertainties tend to exhibit higher values at the boundaries of objects, such as pedestrians and vehicles, which is consistent with the intuition that the decisions along the boundaries are more challenging.

Predictions at long distances also exhibit a high degree of uncertainty, such as the sky.

V. CONCLUSIONS

In this paper, we propose a novel evidential based Stereo-LiDAR fusion for depth estimation. To better exploit the complementary relationship of LiDAR and stereo modalities, we introduce the MFA module at the feature level for early multimodal fusion. Our network computes pixel-wise uncertainties and use them to combine the depth from stereo matching and that from depth completion to get the final depth map. Compared to previous Stereo-LiDAR fusion techniques, our HCENet is able to produce uncertainties for each modality and achieve trustworthy fusion. Experimental results show that the proposed method outperform state-of-the-art methods.

REFERENCES

- [1] P. Nguyen, K. G. Quach, C. N. Duong, N. Le, X.-B. Nguyen, and K. Luu, "Multi-camera multiple 3d object tracking on the move for autonomous vehicles," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2022, pp. 2569–2578.
- [2] M. Franzius, M. Dunn, N. Einecke, and R. Dirnberger, "Embedded robust visual obstacle detection on autonomous lawn mowers," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2017, pp. 44–52.
- [3] N. Garnett, S. Silberstein, S. Oron, E. Fetaya, U. Verner, A. Ayash, V. Goldner, R. Cohen, K. Horn, and D. Levi, "Real-time category-based and general obstacle detection for autonomous driving," in *Proc. IEEE Int. Conf. Comput. Vis. Workshop (ICCVW)*, 2017, pp. 198–205.
- [4] X. Dong, M. A. Garratt, S. G. Anavatti, and H. A. Abbass, "Towards real-time monocular depth estimation for robotics: A survey," *IEEE Trans. Intell. Transp. Syst.*, vol. 23, no. 10, pp. 16940–16961, 2022.
- [5] J. Uhrig, N. Schneider, L. Schneider, U. Franke, T. Brox, and A. Geiger, "Sparsity invariant cnns," in *Proc. IEEE Int. Conf. 3D Vis. (3DV)*, 2017, pp. 11–20.
- [6] I. Laina, C. Rupprecht, V. Belagiannis, F. Tombari, and N. Navab, "Deeper depth prediction with fully convolutional residual networks," in *Proc. IEEE Int. Conf. 3D Vis. (3DV)*, 2016, pp. 239–248.
- [7] F. Ma and S. Karaman, "Sparse-to-dense: Depth prediction from sparse depth samples and a single image," in *Proc. IEEE Int. Conf. Robot. Autom. (ICRA)*, 2018, pp. 4796–4803.
- [8] G. Xu, X. Wang, X. Ding, and X. Yang, "Iterative geometry encoding volume for stereo matching," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2023, pp. 21 919–21 928.
- [9] Y. Lin, H. Yang, T. Cheng, W. Zhou, and Z. Yin, "Dyspn: Learning dynamic affinity for image-guided depth completion," *IEEE Trans. Circuits Syst. Video Technol.*, pp. 1–1, 2023.
- [10] J.-R. Chang and Y.-S. Chen, "Pyramid stereo matching network," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2018, pp. 5410–5418.
- [11] Z. Liang, Y. Guo, Y. Feng, W. Chen, L. Qiao, L. Zhou, J. Zhang, and H. Liu, "Stereo matching using multi-level cost volume and multi-scale feature constancy," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 43, no. 1, pp. 300–315, 2019.
- [12] S. Cheng, Z. Xu, S. Zhu, Z. Li, L. E. Li, R. Ramamoorthi, and H. Su, "Deep stereo using adaptive thin volume representation with uncertainty awareness," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2020, pp. 2524–2534.
- [13] Y. Mao, Z. Liu, W. Li, Y. Dai, Q. Wang, Y.-T. Kim, and H.-S. Lee, "Uasnet: Uncertainty adaptive sampling network for deep stereo matching," in *Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV)*, 2021, pp. 6311–6319.
- [14] S. M. N. Uddin, S. H. Ahmed, and Y. J. Jung, "Unsupervised deep event stereo for depth estimation," *IEEE Trans. Circuits Syst. Video Technol.*, vol. 32, no. 11, pp. 7489–7504, 2022.
- [15] W. Maddern and P. Newman, "Real-time probabilistic fusion of sparse 3d lidar and dense stereo," in *2016 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, 2016, pp. 2181–2188.
- [16] K. Park, S. Kim, and K. Sohn, "High-precision depth estimation with the 3d lidar and stereo fusion," in *Proc. IEEE Int. Conf. Robot. Autom. (ICRA)*, 2018, pp. 2156–2163.
- [17] X. Cheng, Y. Zhong, M. Harandi, Y. Dai, X. Chang, H. Li, T. Drummond, and Z. Ge, "Hierarchical neural architecture search for deep stereo matching," *Proc. Adv. Neural Inf. Process. Syst. (NIPS)*, vol. 33, pp. 22 158–22 169, 2020.
- [18] Y. Zhang and T. Funkhouser, "Deep depth completion of a single rgb-d image," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2018, pp. 175–185.
- [19] K. Zeng, Y. Wang, J. Mao, C. Liu, W. Peng, and Y. Yang, "Deep stereo matching with hysteresis attention and supervised cost volume construction," *IEEE Trans. Image Process.*, vol. 31, pp. 812–822, 2021.
- [20] Z. Li, W. Zuo, Z. Wang, and L. Zhang, "Confidence-based large-scale dense multi-view stereo," *IEEE Trans. Image Process.*, vol. 29, pp. 7176–7191, 2020.
- [21] H.-C. Yang, P.-H. Chen, K.-W. Chen, C.-Y. Lee, and Y.-S. Chen, "Fade: Feature aggregation for depth estimation with multi-view stereo," *IEEE Trans. Image Process.*, vol. 29, pp. 6590–6600, 2020.
- [22] J. Li, Z. Lu, Y. Wang, J. Xiao, and Y. Wang, "Nr-mvsnets: Learning multi-view stereo based on normal consistency and depth refinement," *IEEE Trans. Image Process.*, 2023.
- [23] Z. Lu, J. Wang, Z. Li, S. Chen, and F. Wu, "A resource-efficient pipelined architecture for real-time semi-global stereo matching," *IEEE Trans. Circuits Syst. Video Technol.*, vol. 32, no. 2, pp. 660–673, 2022.
- [24] Y. Lee and H. Kim, "A high-throughput depth estimation processor for accurate semiglobal stereo matching using pipelined inter-pixel aggregation," *IEEE Trans. Circuits Syst. Video Technol.*, vol. 32, no. 1, pp. 411–422, 2022.
- [25] A. Kendall, H. Martirosyan, S. Dasgupta, P. Henry, R. Kennedy, A. Bachrach, and A. Bry, "End-to-end learning of geometry and context for deep stereo regression," in *Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV)*, 2017, pp. 66–75.
- [26] X. Guo, K. Yang, W. Yang, X. Wang, and H. Li, "Group-wise correlation stereo network," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2019, pp. 3273–3282.
- [27] H. Xu and J. Zhang, "Aanet: Adaptive aggregation network for efficient stereo matching," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2020, pp. 1959–1968.
- [28] Z. Li, X. Liu, N. Drenkow, A. Ding, F. X. Creighton, R. H. Taylor, and M. Unberath, "Revisiting stereo depth estimation from a sequence-to-sequence perspective with transformers," in *Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV)*, 2021, pp. 6197–6206.
- [29] Z. Shen, Y. Dai, X. Song, Z. Rao, D. Zhou, and L. Zhang, "Pcw-net: Pyramid combination and warping cost volume for stereo matching," in *Proc. Eur. Conf. Comput. Vis. (ECCV)*, 2022, pp. 280–297.
- [30] J. Qiu, Z. Cui, Y. Zhang, X. Zhang, S. Liu, B. Zeng, and M. Pollefeys, "DeepLidar: Deep surface normal guided depth prediction for outdoor scene from sparse lidar data and single color image," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2019, pp. 3313–3322.
- [31] Y. Zhong, C.-Y. Wu, S. You, and U. Neumann, "Deep rgb-d canonical correlation analysis for sparse depth completion," *Proc. Adv. Neural Inf. Process. Syst. (NIPS)*, vol. 32, 2019.
- [32] Z. Yan, K. Wang, X. Li, Z. Zhang, J. Li, and J. Yang, "Rignet: Repetitive image guided network for depth completion," in *Proc. Eur. Conf. Comput. Vis. (ECCV)*, 2022, pp. 214–230.
- [33] X. Cheng, P. Wang, and R. Yang, "Depth estimation via affinity learned with convolutional spatial propagation network," in *Proc. Eur. Conf. Comput. Vis. (ECCV)*, 2018, pp. 103–119.
- [34] G. Chen, J. Lin, and H. Qin, "Uamd-net: A unified adaptive multimodal neural network for dense depth completion," *IEEE Trans. Circuits Syst. Video Technol.*, vol. 33, no. 10, pp. 5406–5419, 2023.
- [35] Y. Wang, Y. Mao, Q. Liu, and Y. Dai, "Decomposed guided dynamic filters for efficient rgb-guided depth completion," *IEEE Trans. Circuits Syst. Video Technol.*, pp. 1–1, 2023.
- [36] Y. Yang, A. Wong, and S. Soatto, "Dense depth posterior (ddp) from single image and sparse range," in *Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV)*, 2019, pp. 3353–3362.
- [37] B. Forkel and H.-J. Wuensche, "Lidar-sgm: Semi-global matching on lidar point clouds and their cost-based fusion into stereo matching," in *2023 IEEE International Conference on Robotics and Automation (ICRA)*, 2023, pp. 2841–2847.
- [38] M. Hu, S. Wang, B. Li, S. Ning, L. Fan, and X. Gong, "Penet: Towards precise and efficient image guided depth completion," in *Proc. Int. Conf. Robot. Autom. (ICRA)*, 2021, pp. 13 656–13 662.
- [39] S. Zhao, M. Gong, H. Fu, and D. Tao, "Adaptive context-aware multimodal network for depth completion," *IEEE Trans. Image Process.*, vol. 30, pp. 5264–5276, 2021.

- [40] K. Rho, J. Ha, and Y. Kim, "Guideformer: Transformers for image guided depth completion," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2022, pp. 6250–6259.
- [41] Y. Zhang, X. Guo, M. Poggi, Z. Zhu, G. Huang, and S. Mattoccia, "Completionformer: Depth completion with convolutions and vision transformers," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2023, pp. 18 527–18 536.
- [42] B. Lakshminarayanan, A. Pritzel, and C. Blundell, "Simple and scalable predictive uncertainty estimation using deep ensembles," *Proc. Adv. Neural Inf. Process. Syst. (NIPS)*, vol. 30, 2017.
- [43] Y. Gal and Z. Ghahramani, "Dropout as a bayesian approximation: Representing model uncertainty in deep learning," in *Proc. 15th Int. Conf. Mach. Learn. (ICML)*. PMLR, 2016, pp. 1050–1059.
- [44] Z. Han, C. Zhang, H. Fu, and J. T. Zhou, "Trusted multi-view classification with dynamic evidential fusion," *IEEE Trans. Pattern Anal. Mach. Intell.*, 2022.
- [45] W. Van Gansbeke, D. Neven, B. De Brabandere, and L. Van Gool, "Sparse and noisy lidar completion with rgb guidance and uncertainty," in *Proc. 16th Int. Conf. Mach. Vis. Appl. (MVA)*, 2019, pp. 1–6.
- [46] Z. Shen, Y. Dai, and Z. Rao, "Cfnet: Cascade and fused cost volume for robust stereo matching," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, June 2021, pp. 13 906–13 915.
- [47] A. Amini, W. Schwarting, A. Soleimany, and D. Rus, "Deep evidential regression," *Proc. Adv. Neural Inf. Process. Syst. (NIPS)*, vol. 33, pp. 14 927–14 937, 2020.
- [48] C. Wang, X. Wang, J. Zhang, L. Zhang, X. Bai, X. Ning, J. Zhou, and E. Hancock, "Uncertainty estimation for stereo matching based on evidential deep learning," *Pattern Recognition*, vol. 124, p. 108498, 2022.
- [49] J. Lou, W. Liu, Z. Chen, F. Liu, and J. Cheng, "Elfnet: Evidential local-global fusion for stereo matching," *arXiv preprint arXiv:2308.00728*, 2023.
- [50] J. Zhang, M. S. Ramanagopal, R. Vasudevan, and M. Johnson-Roberson, "Listereo: Generate dense depth maps from lidar and stereo imagery," in *Proc. IEEE Int. Conf. Robot. Autom. (ICRA)*, 2020, pp. 7829–7836.
- [51] K. Park, S. Kim, and K. Sohn, "High-precision depth estimation using uncalibrated lidar and stereo fusion," *IEEE Trans. Intell. Transp. Syst.*, vol. 21, no. 1, pp. 321–335, 2019.
- [52] T.-H. Wang, H.-N. Hu, C. H. Lin, Y.-H. Tsai, W.-C. Chiu, and M. Sun, "3d lidar and stereo fusion using stereo matching network with conditional cost volume normalization," in *Proc. IEEE/RSJ Int. Conf. Intell. Robots Syst. (IROS)*, 2019, pp. 5895–5902.
- [53] J. Choe, K. Joo, T. Imtiaz, and I. S. Kweon, "Volumetric propagation network: Stereo-lidar fusion for long-range depth estimation," *IEEE Robot. Automat. Lett.*, vol. 6, no. 3, pp. 4672–4679, 2021.
- [54] X. Cheng, Y. Zhong, Y. Dai, P. Ji, and H. Li, "Noise-aware unsupervised deep lidar-stereo fusion," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2019, pp. 6332–6341.
- [55] Z. Xu, Y. Li, S. Zhu, and Y. Sun, "Expanding sparse lidar depth and guiding stereo matching for robust dense depth estimation," *IEEE Robot. Automat. Lett.*, vol. 8, no. 3, pp. 1479–1486, 2023.
- [56] Q. He, Z. Wang, H. Zeng, Y. Zeng, Y. Liu, S. Liu, and B. Zeng, "Stereo rgb and deeper lidar-based network for 3d object detection in autonomous driving," *IEEE Trans. Intell. Transp. Syst.*, vol. 24, no. 1, pp. 152–162, 2022.
- [57] X. Chen, K.-Y. Lin, J. Wang, W. Wu, C. Qian, H. Li, and G. Zeng, "Bi-directional cross-modality feature propagation with separation-and-aggregation gate for rgb-d semantic segmentation," in *Proc. Eur. Conf. Comput. Vis. (ECCV)*, 2020, pp. 561–577.
- [58] H. Ma, Z. Han, C. Zhang, H. Fu, J. T. Zhou, and Q. Hu, "Trustworthy multimodal regression with mixture of normal-inverse gamma distributions," *Proc. Adv. Neural Inf. Process. Syst. (NIPS)*, vol. 34, pp. 6881–6893, 2021.
- [59] Y. Cabon, N. Murray, and M. Humenberger, "Virtual kitti 2," *arXiv preprint arXiv:2001.10773*, 2020.
- [60] A. Gaidon, Q. Wang, Y. Cabon, and E. Vig, "Virtual worlds as proxy for multi-object tracking analysis," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2016, pp. 4340–4349.
- [61] Y. Zhang, L. Wang, K. Li, Z. Fu, and Y. Guo, "Slfnet: A stereo and lidar fusion network for depth completion," *IEEE Robot. Automat. Lett.*, vol. 7, no. 4, pp. 10 605–10 612, 2022.
- [62] J. Park, K. Joo, Z. Hu, C.-K. Liu, and I. So Kweon, "Non-local spatial propagation network for depth completion," in *Proc. Eur. Conf. Comput. Vis. (ECCV)*, 2020, pp. 120–136.
- [63] J. Tang, F.-P. Tian, W. Feng, J. Li, and P. Tan, "Learning guided convolutional network for depth completion," *IEEE Trans. Image Process.*, vol. 30, pp. 1116–1129, 2020.