

Any to Full: Prompting Depth Anything for Depth Completion in One Stage

Zhiyuan Zhou¹, Ruofeng Liu² , Taichi Liu¹, Weijian Zuo³, Shanshan Wang¹,
Zhiqing Hong⁴, and Desheng Zhang¹

¹ Rutgers University, USA {zhiyuan.z, taichi.liu, shanshan.wang}@rutgers.edu,
desheng@cs.rutgers.edu

² Michigan State University, USA liuruofe@msu.edu

³ JD Logistics, China zuoweijian1@jd.com

⁴ HKUST (Guangzhou), China zhiqinghong@hkust-gz.edu.cn

Abstract. Accurate, dense depth estimation is crucial for robotic perception, but commodity sensors often yield sparse or incomplete measurements due to hardware limitations. Existing RGBD-fused depth completion methods learn priors jointly conditioned on training RGB distribution and specific depth patterns, limiting domain generalization and robustness to various depth patterns. Recent efforts leverage monocular depth estimation (MDE) models to introduce domain-general geometric priors, but current two-stage integration strategies relying on explicit relative-to-metric alignment incur additional computation and introduce structured distortions. To this end, we present Any2Full, a one-stage, domain-general, and pattern-agnostic framework that reformulates completion as a scale-prompting adaptation of a pretrained MDE model. To address varying depth sparsity levels and irregular spatial distributions, we design a Scale-Aware Prompt Encoder. It distills scale cues from sparse inputs into unified scale prompts, guiding the MDE model toward globally scale-consistent predictions while preserving its geometric priors. Extensive experiments demonstrate that Any2Full achieves superior robustness and efficiency. It outperforms OMNI-DC by 32.2% in average AbsREL and delivers a 1.4 $\times$ speedup over PriorDA with the same MDE backbone, establishing a new paradigm for universal depth completion. Codes and checkpoints are available at <https://github.com/zhiyuandaily/Any2Full>.

1 Introduction

Accurate and fine-grained depth information is essential for robotic perception, enabling navigation [36, 53], manipulation [14, 35, 50], and scene understanding [2, 7, 37]. However, commodity depth sensors (e.g., LiDAR [17], ToF [16], or structured light cameras [47]) often yield sparse or incomplete depth maps due to limitations in resolution, range, and light interaction, such as reflection or absorption [59]. Consequently, depth completion has emerged as a fundamental

 Corresponding author.

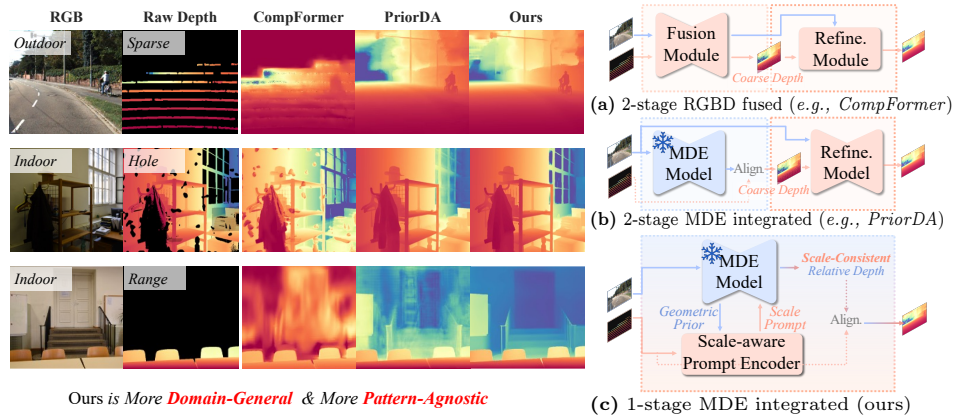


Fig. 1: Depth completion recovers dense depth maps from raw measurements and RGB guidance. (a) Traditional two-stage methods predict coarse depth to bridge the gap between sparse input and dense output. (b) Recent approaches integrate monocular depth estimation (MDE) to generate relative depth and explicitly align it with sparse depth, disrupting MDE’s geometric priors. (c) Our framework employs a lightweight prompting mechanism that tightly integrates MDE’s geometric priors, achieving domain-general and pattern-agnostic depth completion in one stage.

task that aims to recover dense metric depth maps from raw depth measurements and their corresponding RGB images.

Most existing depth completion methods, as illustrated in Fig. 1(a), learn geometric priors from RGB images to guide dense depth prediction from sparse inputs [10, 23, 40, 45, 51, 52, 72]. However, as shown in Fig. 1(left), these RGBD fused frameworks, such as CompFormer [72], learn priors jointly conditioned on RGB distribution and specific depth patterns observed during training, leading to two commonly observed limitations: (1) **Domain Specificity**, where performance degrades under visual domain shifts such as lighting, texture, or scene variations. (2) **Depth Pattern Sensitivity**, where performance degrades with changing raw depth patterns caused by heterogeneous depth sensors, such as variations in density, missing regions, or limitations in sensing range.

Recent monocular depth estimation (MDE) models [24, 68] show remarkable domain generalization, inspiring recent efforts [8, 56, 65] to leverage their geometric priors for depth completion. Nevertheless, current integration strategies remain suboptimal. Test-time adaptation methods [56] are too computationally expensive for real-time use, while direct feed-forward methods [39, 65] typically follow a traditional two-stage design (Fig. 1(a), (b)): an initial stage predicts coarse depth to bridge the gap between sparse to dense output, followed by a refinement stage to recover structural details. In these methods, the RGBD fusion in the first stage is replaced by explicit alignment between relative depth from MDE and sparse metric inputs to produce a coarse metric map.

However, despite effectively addressing the input depth sparsity, this two-stage paradigm not only incurs extra computational overhead but also suffers

from MDE’s inherent scale inconsistency. As analyzed in Supp. A.1, relative depth predictions often exhibit non-linear distortions that require spatially varying scale factors to align with metric ground truth [19]. Explicitly aligning such inconsistent depth maps introduces structured distortions and depth-pattern-specific biases [64], ultimately resulting in noticeable artifacts on unseen depth patterns (i.e. Range), as illustrated by the PriorDA results in Fig. 1 (left). These observations raise a fundamental question: *How can we bypass generation of intermediate coarse depth and instead seamlessly integrate MDE’s geometric priors in one stage, achieving domain-general and pattern-agnostic depth completion?*

To answer this question, we propose Any2Full (Fig. 1 (c)), a one-stage framework that reformulates depth completion as a scale-prompting adaptation of MDE. It distills scale cues (i.e., inter-point scale ratios) from sparse inputs to prompt the pretrained MDE model toward scale-consistent predictions while preserving its robust domain generalization. Here, “**scale-consistent**” means that the predicted relative depth can be aligned to metric depth using a global scale and bias across the scene. However, constructing effective scale prompts from sparse depth is challenging, as sparse inputs not only exhibit varying sparsity levels but also follow irregular spatial distributions (e.g., random missing regions). Such pattern heterogeneity leads to unstable training and often triggers pattern-specific overfitting. To address these issues, we reconstruct the structural context of sparse depth input into unified scale prompt representation under the guidance of MDE’s geometric priors, preserving essential scale cues while decoupling the scale-prompting process from specific depth patterns.

Specifically, we design a Scale-Aware Prompt Encoder with two hierarchical modules. First, the Local Enrichment module couples scale cues from sparse depth with MDE’s dense geometric context, producing local scale-aware features that are robust to varying sparsity levels. Second, the Global Propagation module further propagates these features across the scene through MDE geometry-aware attention, addressing irregular spatial distributions and establishing globally consistent scale-aware features. Finally, the resulting features form a unified scale prompt that modulates MDE features, enabling one-stage, domain-general, and pattern-agnostic depth completion.

In summary, this work has the following contributions:

- We introduce **Any2Full**, a one-stage framework that completes *any* incomplete depth input into a *full* dense depth map. By reformulating depth completion as a scale-prompting adaptation of monocular depth estimation (MDE), it unleashes the power of pretrained MDE models for domain-general and pattern-agnostic depth completion.
- We design a Scale-Aware Prompt Encoder that hierarchically transforms sparse and irregular scale cues into globally consistent, pattern-invariant scale-aware features under MDE geometric guidance, enabling robust scale prompting with minimal additional computation.
- Extensive experiments across diverse domains and depth patterns demonstrate that Any2Full outperforms SOTA (OMNI-DC) by 32.2% in average AbsREL. Compared to the competitive PriorDA using the same MDE back-

bone, Any2Full achieves a $1.4\times$ speedup with superior precision. These results highlight the efficiency of our one-stage scale prompting framework, while its real-world deployment in robotic warehouse grasping further demonstrates its practical value.

2 Related Work

2.1 Monocular Depth Estimation

Monocular depth estimation (MDE) [15, 44] aims to infer the 3D scene structure from a single RGB image. It is an ill-posed inverse problem due to inherent scale ambiguity, which means that many distinct 3D scenes can produce identical image projections. Nonetheless, relative depth relationships can be inferred from visual cues such as texture gradients, object sizes, and occlusions [1].

Inspired by the success of foundation models in natural language processing [5, 12], MDE has evolved toward strong zero-shot generalization [3, 4, 18, 24, 38, 44, 61, 62, 67]. Early efforts like MiDaS [44] leveraged mixed-dataset training, while Depth Anything [68] utilizes massive unlabeled images to learn robust relative geometric priors. In parallel, Marigold [24] converts pre-trained latent diffusion models into conditional depth estimators to achieve superior generalization. While the above MDE models are limited to relative depth prediction due to scale ambiguity, some recent methods [3, 69, 73] like UniDepth [42, 43] and WorDepth [71] incorporate camera or object priors to enable monocular metric depth estimation. Nevertheless, the scale ambiguity remains fundamentally unsolved in the absence of reliable metric measurements.

2.2 Depth Completion

Depth completion [34] aims to recover dense metric depth from sparse inputs and RGB images. Conventional methods adopt a two-stage pipeline [9–11, 30, 31, 40, 51, 60, 72]: an encoder–decoder network first fuses depth and RGB to predict a coarse dense depth, which is then refined by spatial propagation modules guided by RGB. Although effective on benchmarks, these models are highly domain- and pattern-specific, struggling under appearance shifts (lighting, texture, scene type) or heterogeneous real-world depth patterns such as sparsity, missing regions [60], and range limitations caused by sensor constraints.

In contrast to data-centric paradigms that train completion models from scratch on massive datasets [33, 57, 58, 74], recent efforts integrate MDE into depth completion to improve cross-domain robustness by combining the rich geometric priors of pretrained MDE with metric cues from sparse depth. Two families have emerged. *Test-time adaptation methods* [19–21, 56], such as Marigold-DC [56] and TestPromptDC [20], leverage sparse depth as supervision to iteratively optimize model outputs during inference, achieving high metric accuracy while preserving the MDE pipeline. However, such approaches incur heavy inference costs, limiting their real-time robotic applications. In contrast, *Direct*

feed-forward methods [39, 65, 70] prioritize efficiency but face the challenge of adapting relative depth models to metric prediction. PriorDA [65] adopts a two-stage pipeline: it first explicitly aligns the frozen MDE’s relative output with sparse measurements to obtain a coarse dense map, then refines it using a secondary RGB-guided model. However, we argue that such *output-level fusion* distorts MDE’s geometric priors because the relative-to-metric mapping is often spatially inconsistent. This introduces structural artifacts that the refinement stage must laboriously correct, ultimately making the pipeline heavier and more sensitive to unseen depth patterns (detailed two-stage limitation analysis in Supp. A.2). Meanwhile, PromptDA [29], designed for depth super-resolution, does not transfer well to depth completion due to its specific depth pattern.

To overcome these limitations, we rethink depth completion as a one-stage scale-prompting adaptation of MDE and design a scale-aware prompt encoder that constructs unified scale prompts for varying depth patterns, effectively integrating MDE’s geometric priors with sparse depth cues to achieve domain-general and pattern-agnostic performance.

3 Any2Full

In this section, we present Any2Full, a one-stage depth completion framework that is domain-general and depth-pattern-agnostic. Sec. 3.1 introduces our scale-prompting formulation, which adapts a pretrained monocular depth estimation (MDE) model for depth completion by injecting scale cues from sparse inputs and leveraging its domain-general geometric priors. Sec. 3.2 details the scale prompting pipeline, centered around a scale-aware prompt encoder that constructs unified scale prompts for diverse depth patterns. Finally, Sec. 3.3 describes training objectives and implementation details.

3.1 Formulation

Depth completion aims to recover a dense metric depth map $\mathbf{D}_f \in \mathbb{R}^{H \times W}$ from an RGB image $\mathbf{I} \in \mathbb{R}^{H \times W \times 3}$ and its corresponding sparse measurements $\mathbf{D}_s$. Traditional approaches directly learn a mapping $\hat{\mathbf{D}}_f = \mathcal{F}_\theta(\mathbf{I}, \mathbf{D}_s)$ by fusing RGB and sparse depth, which often results in domain-specific representations.

To overcome this limitation, we **reformulate** depth completion as a scale-prompting adaptation of pretrained MDE models. Fig. 2 (a) illustrates our one-stage framework Any2Full, where depth inputs are encoded into scale prompts that modulate pretrained MDE features to produce *scale-consistent* relative depth, which can be converted to metric depth using a single global scale and bias (as analyzed in Supp. A.1). This formulation enables the model to inherit the domain-general geometric priors of a pretrained MDE backbone.

Formally, given a pretrained MDE model $\mathcal{M}$, an RGB image $\mathbf{I}$ and a raw depth map $\mathbf{D}_s$, our goal is to estimate a dense metric depth $\hat{\mathbf{D}}_f$. To bridge the gap between sparse metric observations and the MDE’s relative representation, we first normalize $\mathbf{D}_s$ into $\hat{\mathbf{D}}_s$ by removing global scale and bias while preserving

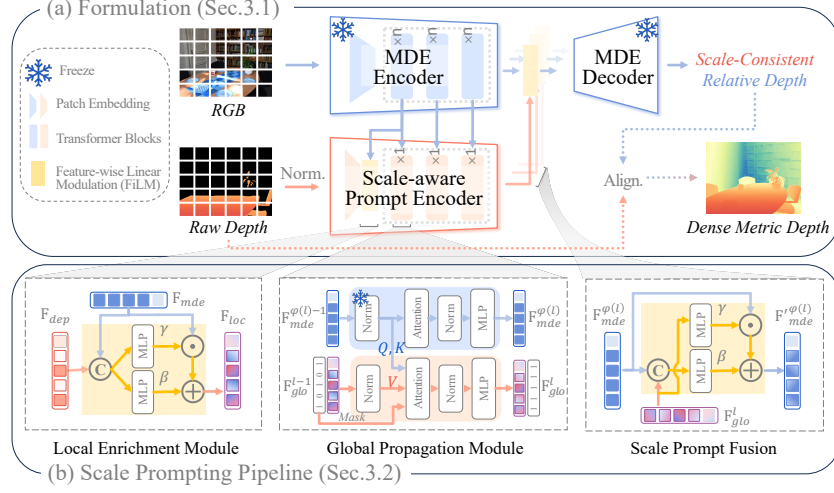


Fig. 2: Overview of Any2Full. (a) Our framework reformulates depth completion as a scale-prompting adaptation of a pretrained MDE model. The normalized raw depth is encoded into scale prompts to modulate MDE features for scale-consistent relative depth prediction, followed by a non-parametric least-squares fit to recover the dense metric depth. (b) The Scale-aware Prompt Encoder (SAPE) transforms raw depth into unified scale prompts through two hierarchical modules: *Local Enrichment* anchors scale cues into the MDE latent space; and *Global Propagation* leverages MDE geometric features as guidance to diffuse scale cues across patches. Finally, *Scale Prompt Fusion* injects the resulting prompts into the MDE decoder for final prediction.

inter-point scale ratios as scale cues (details in Supp. C.1). We then employ a **Scale-aware Prompt Encoder (SAPE)**, denoted as $\mathcal{G}$ to convert $\tilde{\mathbf{D}}_s$ into a scale prompt that modulates the features of $\mathcal{M}$ to produce a scale-consistent relative prediction

$$\hat{\tilde{\mathbf{D}}}_f = \mathcal{M}(\mathbf{I} \mid \mathcal{G}(\tilde{\mathbf{D}}_s)). \quad (1)$$

Finally, the relative prediction $\hat{\tilde{\mathbf{D}}}_f$ is aligned with the raw metric depth $\mathbf{D}_s$ via non-parametric least-squares fit [44] to recover the dense metric depth $\hat{\mathbf{D}}_f$.

One-stage inference. We define one-stage as a single forward inference process without intermediate depth prediction or auxiliary refinement networks. Our framework satisfies this definition because the final metric depth is obtained via a closed-form alignment applied to the predicted relative depth, which introduces no additional learnable modules.

Overall, our formulation bridges depth completion and monocular depth estimation within a one stage scale-prompting framework preserving MDE’s geometric priors while injecting scale cues for domain-general depth completion.

3.2 Scale Prompting Pipeline

The core challenge in designing the SAPE $\mathcal{G}$ within this scale prompting pipeline lies in handling varying depth patterns, which exhibit different sparsity levels and irregular spatial distributions (i.e., random missing regions), leading to unstable training and triggering pattern-specific overfitting.

To address this, SAPE progressively reconstructs the structural context of sparse depth input into a unified scale prompt representation under the guidance of MDE’s geometric priors, with two hierarchical modules: (1) a local enrichment module that produces local scale-aware features robust to sparsity variations, and (2) a global propagation module that handles irregular spatial distributions and establishes a globally consistent representation as scale prompt. Then, a scale prompt fusion module uses the resulting scale prompt to modulate MDE features for scale-consistent prediction, forming our scale prompting pipeline. Fig. 2 (b) details the components.

Throughout, we build upon a ViT-style [13] MDE backbone [68] and the SAPE follows ViT encoder with a lightweight design. Given the sparse depth input $\tilde{\mathbf{D}}_s$, the SAPE first embeds it into patch-level depth features $\mathbf{F}_{dep}$ (see Supp. C.2 for details). These features are subsequently encoded into unified scale prompts via the following two hierarchical modules:

Local Enrichment. To achieve robustness under varying sparsity levels, the Local Enrichment module distills scale cues from $\mathbf{F}_{dep}$ and couples them with the dense geometric features $\mathbf{F}_{mde}$ from the MDE backbone, producing local scale-aware features $\mathbf{F}_{loc}$ that are insensitive to sparsity variations.

In detail, each patch-level depth feature $\mathbf{f}_{dep,i}$ interacts with its corresponding MDE feature $\mathbf{f}_{mde,i}$ through generalized Feature-wise Linear Modulation (FiLM) [41] mechanism to anchor scale cues within the MDE latent space:

$$\mathbf{f}_{loc,i} = \gamma(\mathbf{f}_{dep,i}, \mathbf{f}_{mde,i}) \odot \mathbf{f}_{mde,i} + \beta(\mathbf{f}_{dep,i}, \mathbf{f}_{mde,i}), \quad (2)$$

where $\odot$ denotes element-wise multiplication, $\gamma(\cdot)$ and $\beta(\cdot)$ are modulation parameters predicted by a lightweight Multi-Layer Perceptron (MLP).

Global Propagation. While the Local Enrichment module anchors scale cues within individual patches, these cues remain spatially fragmented and lack global consistency. To bridge this gap, the Global Propagation module performs MDE geometry-guided diffusion.

Specifically, we treat the local scale-aware features $\mathbf{F}_{loc}$ as the initial state $\mathbf{F}_{glo}^0 = \mathbf{F}_{loc}$, and update them through L geometry-guided Transformer blocks, where L matches the MDE encoder groups (and corresponding decoder levels):

$$\mathbf{F}_{glo}^l = \text{TransformerBlock}(\mathbf{F}_{glo}^{l-1}, \mathbf{F}_{mde}^{\phi(l)-1}), \quad l = 1:L. \quad (3)$$

where $\mathbf{F}_{mde}^{\phi(l)-1}$ denotes the geometry features extracted from the last block of the l -th MDE encoder group, and $\phi(l) = l \cdot n$ is the index mapping function (assuming each group contains n blocks). Different from standard cross-attention mechanisms, our module computes attention weights purely from MDE geometry

(Query and Key from $\mathbf{F}_{mde}$), using $\mathbf{F}_{glo}$ only as the Value. This ensures that scale cues are diffused along geometric structures defined by the RGB-based MDE features rather than being biased by the irregular sampling patterns of the sparse depth. The resulting globally consistent scale-aware features $\mathbf{F}_{glo}^l$ provide a unified scale prompt for diverse depth patterns that benefit subsequent prompt-based fusion in a pattern-agnostic manner.

Scale Prompt Fusion. The set of scale prompts $\{\mathbf{F}_{glo}^1, \dots, \mathbf{F}_{glo}^L\}$ is injected into the MDE decoder for scale-consistent prediction through hierarchical Feature-wise Linear Modulation (FiLM). Specifically, at each decoder level l , the modulation is applied to $\mathbf{F}_{mde}^{\phi(l)}$, which is the intermediate MDE feature within the decoder path prior to its fusion with other level features:

$$\mathbf{F}_{mde}^{\phi(l)} = \gamma^l(\mathbf{F}_{glo}^l) \mathbf{F}_{mde}^{\phi(l)} + \beta^l(\mathbf{F}_{glo}^l, \mathbf{F}_{mde}^{\phi(l)}), \quad l = 1:L. \quad (4)$$

$\gamma^l(\cdot)$ and $\beta^l(\cdot)$ are predicted by lightweight MLP specialized for each level. Each modulation is applied token-wise within its level. This multi-level fusion injects scale prompt at multiple semantic levels, refining relative depth features toward scale-consistent predictions while preserving the domain-general geometric priors from the MDE backbone.

Efficiency and Lightweight Design. The Scale-aware Prompt Encoder is designed for minimal overhead while ensuring effective guidance. To achieve this, the number of propagation blocks is specifically aligned with the number of hierarchical levels in the MDE decoder, ensuring that each scale prompt semantically corresponds to a specific decoder level without redundant computation. Moreover, propagation and prompting are omitted for the lowest-level encoder-decoder layers, where features primarily capture local textures rather than the structural geometric priors necessary for scale propagation.⁵ To prevent early-propagation loss of sparse scale cues, we adopt masked attention in the first transformer block, constraining scale information spread from valid depth locations to geometrically coherent regions.

Design Analysis. *Pattern-agnostic.* Scale-aware Prompt Encoder adapts to diverse depth patterns. The depth input provides only scale cues at valid points, while enrichment and propagation rely on MDE geometry, ensuring invariance to sparsity levels and spatial distributions. *Domain-general.* Unlike conventional prompting methods [22, 28] that inject task-specific tokens at the input level to condition model behavior, our approach applies lightweight feature-level modulation via FiLM using a scale prompt. This design minimally alters the original MDE representations, injecting only scale-related cues to enhance scale consistency while preserving domain-general priors.

⁵ For example, with a ViT-L backbone [68] providing intermediate features at layers {5, 11, 17, 23}, we enrich local features using layer 10, perform three propagation steps guided by layers {10, 16, 22}, and execute prompt fusion at levels corresponding to layers {11, 17, 23}.

3.3 Training

Training Data. Since Any2Full requires RGB-D pairs with accurate metric ground truth, we train our model on high-quality synthetic datasets. To encourage robustness to varying depth patterns, we adopt two depth sampling strategies: (1) Random Sampling, which randomly selects depth points to produce sparse depth maps with varying density; and (2) Hole Sampling, following the method of [60], produces depth inputs with large contiguous missing regions. During training, we randomly select one of the two strategies for each sample to ensure pattern diversity.

Training Objective. We supervise the predicted dense depth using a combination of losses that enforce global scale consistency, local structure sharpness, and sparse input alignment. Following prior MDE works [44, 59, 67], we adopt a scale- and shift-invariant loss $\mathcal{L}_{\text{ssi}}$ for global consistency and a gradient matching loss $\mathcal{L}_{\text{gm}}$ for edge preservation. To ensure consistency with sparse measurements, we additionally apply $\mathcal{L}_{\text{ssi}}$ on valid depth anchors, denoted as $\mathcal{L}_{\text{anchor}}$. We also include a Relative-Structure SSIM loss $\mathcal{L}_{r\text{-ssim}}$ [19] to regularize structural similarity between prediction and ground truth. The overall training objective is a weighted sum of these losses (see Supp. C.3 for details).

Implementation Details. We utilize *Depth Anything v2* [68] as the MDE backbone, initialized from its released relative-depth model. We freeze the MDE backbone and train the scale-aware prompt encoder on a subset of the Depth Anything v2 synthetic dataset, including 60K indoor images from *Hypersim* [46], 10K outdoor images sampled from *VKITTI2* [6], and 15K images sampled from *TartanAir* [63] covering both indoor and outdoor scenes. The sampling is conducted to maintain dataset balance and diversity while considering computational constraints. The model is trained for 224K steps with a 10K-step warm-up phase and a cosine scheduler [32]. We use a batch size of 16 on four NPUs (equivalent to high-end GPU clusters). The optimizer is Adam [25] with a learning rate of $5\text{e-}5$ for the scale-aware prompt encoder.

4 Experiments

4.1 Experimental Setup

Public Benchmarks. To assess the zero-shot generalization ability of Any2Full across visual domains, we evaluate on six unseen public datasets spanning indoor and outdoor scenes with diverse visual conditions and sensor modalities, forming a comprehensive benchmark for cross-domain depth completion. **NYU-Depth V2** [49] contains indoor scenes captured by an RGB-D Kinect sensor, and we follow the standard test split of 654 samples at 304×228 resolution. **iBims-1** [26] is a high-quality indoor dataset captured with a laser scanner featuring accurate geometry and an extended depth range of up to 50 m, and we use all 100 test images at 640×480 resolution. **KITTI DC** [54] consists of outdoor driving scenes with sparse LiDAR depth at 1216×352 resolution, and we adopt the standard validation split of 1,000 samples. **DIODE** [55] is a large-scale

indoor-outdoor dataset captured with a FARO laser scanner, providing highly accurate dense depth up to 350 m. We use the official validation split of 771 samples at 640×480 resolution. **ETH3D** [48] is a multi-view stereo benchmark featuring high-resolution indoor and outdoor scenes from DSLR imagery and FARO scans, and we evaluate the High-res DSLR set (11 scenes, 390 images) at 640×480 resolution. **VOID** [66] contains indoor sequences captured by an Intel RealSense D435i sensor, featuring sparse, noisy depth with challenges such as motion blur and low texture. We use the standard test set of 800 samples with three sparsity protocols (1500/500/150) at 640×480 resolution.

Real-World Dataset. Beyond public benchmarks, we further evaluate on a newly collected real-world dataset **Logistic-Black** to assess the practical robustness of our model. This dataset is collected in a robotic warehouse scene using a Vzense DS77C Pro ToF camera. It contains 114 RGB-D pairs cropped to 1448×1048 resolution. Logistic-Black represents a challenging robotic-industrial domain characterized by black packages with top-down camera views. Black packages absorb infrared light, leading to missing regions in ToF depth measurements. To obtain metrically accurate ground truth for evaluation, we replace the package surfaces with reflective materials under the same setup.

Depth Patterns. To comprehensively evaluate the robustness of Any2Full across different depth patterns, we design six representative settings: Hole, Range, Sparse-Random, Sparse-LiDAR, Sparse-SfM, and Mixed. Sparse patterns follow each dataset’s standard protocol. Hole and Range patterns follow and extend robustness protocols in prior studies [60, 65]. Range pattern is unseen for all methods, as none are trained with range-truncated depth inputs. **Hole.** Simulates missing regions commonly observed in indoor RGB-D sensors, where raw depth is dense, but continuous areas are missing due to reflection or absorption of sensor light. Following [59, 60], we adopt highlight masking, black masking, small-block masking, and semantic masking to generate realistic hole distributions. **Range.** Models sensor range limitations [39, 65]. TOF sensors in indoor scenes are typically unreliable at very near or far distances, while outdoor LiDARs are range-limited on reflective surfaces. We retain the central 20–80% of valid depth values for both indoor and outdoor, forming a range-constrained completion task. **Sparse-Random.** This widely used setting randomly samples sparse points from dense ground truth, though it differs from real-world depth distributions. Following [10, 34, 40], we randomly sample 500 points for NYU-Depth V2, 1 000 points for iBims-1 and 0.1% points for DIODE. **Sparse-LiDAR.** This pattern corresponds to outdoor depth captured by LiDAR sensors, we use the original 64-lines LiDAR measurements (64L) as sparse inputs for KITTI DC. **Sparse-SfM.** Following [74], we project COLMAP SfM points into image space to obtain sparse depth maps representative of multi-view reconstruction for ETH3D. **Mixed.** This pattern is specific to the *Logistic-Black* dataset, combining depth holes caused by low-reflectivity black packages and sparsity induced by the ToF sensor’s limited depth resolution ($\sim 20\%$).

Baselines. We compare Any2Full with three method categories: (1) Domain-specific depth completion: CompletionFormer (CompFormer) [72] and Depth-

Table 1: Quantitative comparison on large-scale zero-shot benchmarks (Metrics are AbsREL ($\downarrow$) and RMSE ($\downarrow$); **bold** and underlined denote the best and second-best results). “Any2Full (DA-L/B/S)” denotes our model using Depth Anything-Large/Base/Small backbones. **Our method Any2Full (DA-L) achieves the lowest average rank (2.3) across all scenarios**, consistently performing among the top competitors against strong domain-generalization and pattern-agnostic approaches, demonstrating robust generalization across diverse depth patterns and domains.

Method	Avg Rank	AVG												NYU-Depth V2						IBims-1					
		Hole		Range		Sparse		Hole		Range		Sparse-Random		Hole		Range		Sparse-Random							
		AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE						
		↓																							
CompFormer	7.9	0.316	9.682	0.355	5.499	0.067	8.020	0.342	1.098	0.324	1.050	0.011	0.089	0.371	1.372	0.372	1.511	0.013	0.197						
DepthPrompt [†] *	5.1	0.018	0.816	0.121	3.558	0.098	2.032	0.008	0.097	0.051	0.338	0.014	0.104	0.008	0.145	0.091	0.647	0.017	0.202						
OMNI-DC [‡]	3.4	0.014	0.472	0.082	2.805	0.025	0.648	0.019	0.108	0.073	0.415	0.014	0.112	0.014	0.152	0.076	0.592	0.009	0.148						
Depth Anything [†] *	6.4	0.069	2.361	0.079	2.472	0.074	2.207	0.062	0.313	0.074	0.342	0.061	0.299	0.038	0.329	0.042	0.303	0.038	0.337						
Marigold [†] *	7.5	0.115	2.084	0.294	2.665	0.142	2.230	0.080	0.324	0.099	0.375	0.082	0.320	0.060	0.346	0.075	0.400	0.060	0.354						
PromptDA [†] *	9.1	0.400	6.036	0.644	8.063	0.408	5.643	0.333	1.349	0.588	2.240	0.033	0.224	0.238	1.232	0.533	2.634	0.021	0.191						
Marigold-DC [†] *	5.7	0.024	0.597	0.091	2.995	0.193	2.088	0.017	0.114	0.080	0.412	0.021	0.153	0.016	0.155	0.070	0.535	0.018	0.208						
TestPromptDC [†] *	3.5	0.017	0.445	0.040	1.660	0.161	2.046	0.014	0.070	0.057	0.393	0.018	0.114	0.013	0.120	0.038	0.306	0.015	0.163						
PriorDA [†] *	3.6	0.015	0.655	0.080	2.567	0.028	0.775	0.013	0.102	0.074	0.363	0.016	0.119	0.011	0.141	0.063	0.459	0.011	0.143						
Any2Full(DA-L) [†] *	2.3	0.010	0.459	0.046	2.184	0.026	0.829	0.008	0.067	0.035	0.259	0.017	0.138	0.007	0.110	0.030	0.334	0.013	0.161						
Any2Full(DA-B) [†] *		0.010	0.445	0.046	2.134	0.028	0.822	0.008	0.066	0.035	0.264	0.017	0.131	0.008	0.112	0.031	0.336	0.013	0.167						
Any2Full(DA-S) [†] *		0.011	0.484	0.0549	2.531	0.029	0.877	0.008	0.068	0.041	0.303	0.017	0.126	0.009	0.126	0.039	0.405	0.015	0.174						

Method	DIODE						KITTI DC						ETH3D						Logistic-Black	
	Hole		Range		Sparse-Random		Hole		Range		Sparse-LiDAR		Hole		Range		Sparse-SM		Mixed	
	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE	AbsREL RMSE		
CompFormer	0.353	15.110	0.384	7.617	0.034	11.014	0.163	28.108	0.307	13.370	0.041	18.348	0.350	2.724	0.387	3.946	0.235	10.450	0.011	0.039
DepthPrompt [†] *	0.020	1.091	0.152	4.494	0.029	1.315	0.040	2.428	0.170	9.967	0.098	5.005	0.012	0.319	0.142	2.343	0.332	3.534	0.009	0.025
OMNI-DC [‡]	0.015	0.962	0.100	3.839	0.017	0.915	0.011	0.907	0.079	7.469	0.013	1.230	0.010	0.231	0.081	1.710	0.070	0.835	0.009	0.022
Depth Anything [†] *	0.110	5.061	0.130	5.083	0.110	4.937	0.078	3.507	0.090	4.646	0.085	3.830	0.058	2.594	0.062	1.988	0.075	1.633	0.026	0.059
Marigold [†] *	0.140	2.751	0.140	3.446	0.141	2.770	0.203	6.080	0.206	7.951	0.194	6.192	0.092	0.918	0.950	1.152	0.232	1.513	0.025	0.061
PromptDA [†] *	0.299	7.244	0.567	10.831	0.038	1.691	0.829	16.804	0.916	18.710	0.967	18.960	0.301	3.549	0.614	5.898	0.979	7.149	0.990	1.714
Marigold-DC [†] *	0.034	1.028	0.099	3.755	0.034	1.260	0.041	1.489	0.105	8.516	0.044	1.865	0.013	0.198	0.100	1.755	0.847	6.954	0.011	0.036
TestPromptDC [†] *	0.034	1.073	0.085	3.315	0.033	1.138	0.016	0.836	0.037	3.630	0.020	1.350	0.008	0.118	0.032	0.654	0.719	7.465	0.013	0.040
PriorDA [†] *	0.024	1.218	0.091	3.586	0.022	1.144	0.020	1.623	0.096	6.990	0.022	1.714	0.007	0.190	0.078	1.438	0.070	0.754	0.011	0.037
Any2Full(DA-L) [†] *	0.018	1.164	0.062	3.217	0.022	1.363	0.011	0.784	0.061	6.000	0.017	1.356	0.005	0.173	0.041	1.111	0.063	1.129	0.010	0.029
Any2Full(DA-B) [†] *	0.018	1.116	0.063	3.297	0.022	1.312	0.011	0.771	0.058	5.597	0.017	1.363	0.005	0.161	0.041	1.175	0.069	1.135	0.009	0.028
Any2Full(DA-S) [†] *	0.019	1.257	0.068	3.422	0.023	1.429	0.012	0.794	0.070	7.043	0.018	1.450	0.006	0.178	0.053	1.483	0.070	1.205	0.010	0.030

[†] Domain Generalization method. [‡] Depth pattern-agnostic method. ^{*} MDE-based approach.

Prompting (DepthPrompt) [39], where the abbreviations are used for compactness. (2) Domain-general monocular depth estimation (MDE): Depth Anything v2 [68] and Marigold [24], which predict relative depth from RGB only. (3) Domain-general methods: PromptDA [29] is designed for depth super-resolution. Marigold-DC [56], TestPromptDC [20], PriorDA [65], and OMNI-DC [74] are depth completion methods, where all except OMNI-DC leverage MDE priors.

Evaluation Protocols. All models are evaluated in a zero-shot setting without dataset-specific fine-tuning. For domain-specific methods, we use their released models trained on the NYU-Depth V2 dataset for evaluation. For MDE methods, relative depth predictions are aligned to sparse depth using least-squares fitting to produce metric outputs. All Depth Anything (DA)-based methods use the released DA-v2-Large backbone for consistency. For PromptDA, the raw sparse depth is downsampled via nearest-neighbor interpolation to form dense low-resolution inputs. We compute Absolute Relative Error (AbsREL) and Root Mean Squared Error (RMSE, in meters) [27] on valid pixels under the same cross-domain and depth-pattern settings.

4.2 Comparisons with the State of the Art

Domain Generalization. As shown in Tab. 1, the performance of domain-specific methods degrades sharply when transferred to unseen domains, as CompFormer (NYU→KITTI), revealing that traditional RGBD-fused approaches are overfit to a specific domain. Monocular depth estimation methods like DA and

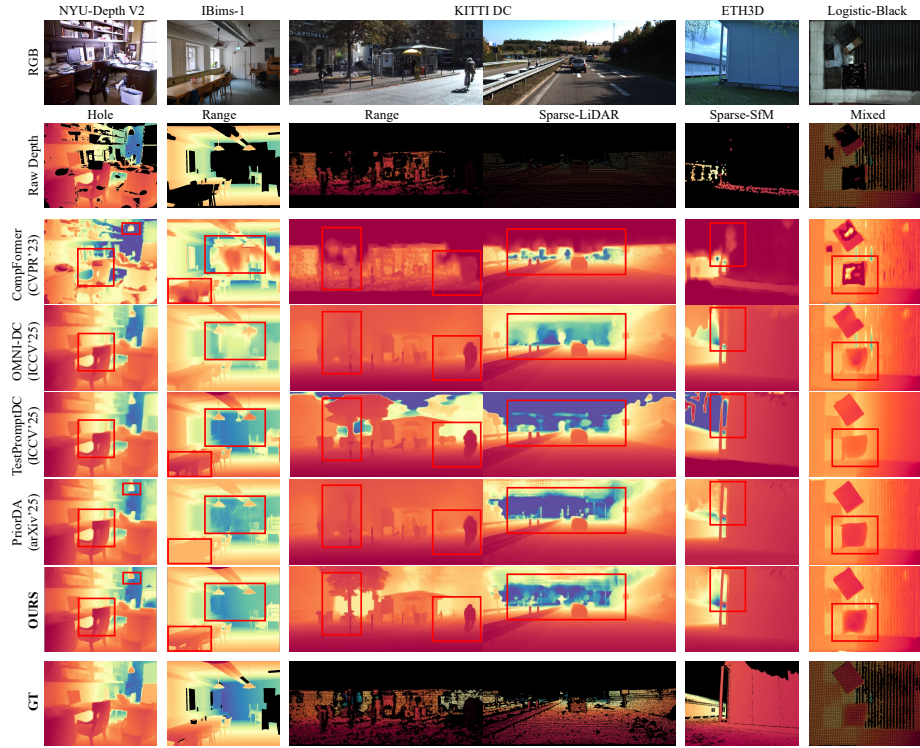


Fig. 3: Qualitative comparison under various depth sampling patterns. Black shows missing depths, and red boxes mark key areas. **Any2Full demonstrates accurate global geometry, structural consistency, and fine-grained detail preservation across all patterns.**

Marigold generalize well for relative depth prediction, but their scale inconsistency causes significant errors when aligning with metric inputs.

OMNI-DC (rank 3.4) improves domain generalization by leveraging large-scale multi-domain training on diverse synthetic datasets combined with scale normalization. Meanwhile, MDE-integrated depth completion methods, including Marigold-DC (rank 5.7), TestPromptDC (rank 3.5), and PriorDA (rank 3.6), enhance cross-domain robustness by incorporating geometric priors from pre-trained MDE models.

In comparison, our one-stage MDE-integrated framework Any2Full further improves both stability and accuracy, achieving the lowest average rank (2.3) across all evaluated scenarios, consistently remaining among the top performers on every dataset. Rather than relying on large-scale training data, our scale-prompting strategy effectively fuses geometric priors from pretrained MDE with scale cues from sparse depth, leading to robust domain-general depth completion.

Depth Pattern Generalization. As shown by the average metrics across various patterns in Tab. 1, existing methods exhibit high sensitivity to depth sam-

Table 2: Model efficiency comparison on IBims-1. Average runtime is computed on 100 640×480 images and hole patterns using a single NVIDIA RTX P40 GPU.

Model	MDE	MDE Params.	Add. Params.	Latency (s)	AbsREL (↓)
Marigold-DC	Marigold	1.29B	0M	112.47	0.016
TestPromptDC	UniDepth	363.21M	0.9M	92.40	0.013
PriorDA	DA-L	335.3M	96.4M	0.68	0.011
Depth Anything	DA-L	335.3M	0M	0.48	0.038
Ours	DA-L	335.3M	60.6M	0.49	0.007
	DA-B	97.5M	34.2M	0.21	0.008
	DA-S	24.8M	8.9M	0.09	0.009

pling patterns compared to Any2Full. Among state-of-the-art pattern-agnostic methods, OMNI-DC and PriorDA appear to overfit to the Sparse setting and generalize poorly to unseen Range patterns, indicating their completion relies heavily on pattern-specific correlations. Meanwhile, test-time adaptation methods like TestPromptDC and Marigold-DC prove unstable in sparse scenarios, suffering significant performance degradation. **Qualitative Results** in Fig. 3 further illustrate these differences. Traditional methods like CompFormer fail entirely. OMNI-DC exhibits noticeable depth bleeding (e.g., in Range and Sparse-SfM). While MDE-based methods benefit from fine-grained details, they introduce distinct artifacts: TestPromptDC produces erroneous depth rings and over-smoothing (i.e., KITTI DC-Range) via its test-time prompt optimization and PriorDA generates artifacts (i.e., IBims-Range) from noisy output-level alignment with MDE predictions. In contrast, Any2Full maintains the fine-grained geometric details provided by the MDE while achieving superior global scale accuracy and structural consistency across all patterns. Furthermore, in the real-world dataset Logistic-Black, Any2Full achieves the most robust visual quality despite OMNI-DC’s slightly lower AbsREL, benefiting downstream robotic tasks. These results confirm that our scale-aware prompt encoder effectively decouples scale cues from sparse depth, enabling pattern-invariant scale prompting.

Model Efficiency. As summarized in Tab. 2, Any2Full achieves state-of-the-art performance while maintaining high computational efficiency. When built on DA-L, our method adds less than 20% additional parameters (60.6M vs. 335.3M) to the MDE backbone, yet reduces AbsREL from 0.038 (Depth Anything) to 0.007. Compared to the competitive PriorDA using the same backbone, Any2Full achieves a **1.4×** speedup (0.49 s vs. 0.68 s) with superior accuracy. Remarkably, the smallest variant of Ours (DA-S) still surpasses all prior depth completion methods, achieving a 0.09 s inference time, about **7×** faster than PriorDA and **1000×** faster than TestPromptDC. This highlights the efficiency of our one-stage scale prompting framework, which avoids the multi-stage refinement overhead of PriorDA or the iterative optimization required by TestPromptDC.

Benefiting from this high efficiency and zero-shot accuracy, **Any2Full has been deployed in a real-world robotic warehouse** cell handling thousands of packages daily. It improved the grasping success rate for challenging black packages from 28% to 91.6% without damage (see Supp. B).

Backbone Consistency. We further train our model with different backbone sizes (Depth Anything-Large, Base, and Small) to examine the effectiveness of

Table 3: Cross-backbone generalization of SAPE. AbsREL ($\downarrow$) is evaluated across six benchmarks and multiple depth patterns. “Any2Full (MoGe-2-B)” denotes our model using MoGe-2-Base as the MDE backbone, while “Any2Full (DA-B)” denotes our model using Depth Anything-Base.

Method	Average			NYU-Depth V2			IBims-1			DIODE			KITTI DC			ETH3D			Logistic-Black	
	Hole	Range	Sparse	Hole	Range	Random	Hole	Range	Random	Hole	Range	Random	Hole	Range	LiDAR	Hole	Range	SF	Mixed	
MoGe-2-B	0.059	0.066	0.084	0.060	0.072	0.061	0.036	0.044	0.036	0.085	0.087	0.086	0.067	0.073	0.079	0.049	0.052	0.158	0.030	
Any2Full (MoGe-2-B)	0.012	0.051	0.034	0.010	0.043	0.018	0.009	0.038	0.014	0.020	0.066	0.025	0.015	0.060	0.021	0.008	0.048	0.090	0.011	
Any2Full (DA-B)	0.010	0.046	0.028	0.008	0.035	0.017	0.008	0.031	0.013	0.018	0.063	0.022	0.011	0.058	0.017	0.005	0.041	0.069	0.009	

Table 4: Robustness comparison on varying depth sparsity levels. We evaluate on the VOID dataset with 1500/500/150 sparse points sampled from a VIO system, and the KITTI dataset with 64/16/4-Line LiDAR scans.

Method	AVG		VOID-1500		VOID-500		VOID-150		KITTI-64L		KITTI-16L		KITTI-4L	
	AbsREL	RMSE	AbsREL	RMSE	AbsREL	RMSE	AbsREL	RMSE	AbsREL	RMSE	AbsREL	RMSE	AbsREL	RMSE
DepthPrompt	0.122	3.187	0.046	0.623	0.080	0.648	0.115	0.755	0.098	5.005	0.106	4.701	0.283	7.391
OMNI-DC	<u>0.032</u>	<u>1.220</u>	0.026	0.555	<u>0.039</u>	<u>0.551</u>	<u>0.057</u>	<u>0.650</u>	0.013	1.230	0.020	<u>1.703</u>	0.039	2.630
TestPromptDC	0.172	1.638	0.039	0.683	0.068	0.655	0.850	3.218	0.020	<u>1.350</u>	0.023	1.699	0.032	2.221
Ours	0.031	1.203	<u>0.030</u>	<u>0.561</u>	0.037	0.538	0.045	0.620	<u>0.017</u>	1.356	<u>0.023</u>	1.760	<u>0.035</u>	<u>2.380</u>

our approach, as shown in Tab. 1. Interestingly, a larger DA-L backbone does not necessarily improve performance, likely due to the coarse ground truth in datasets like NYU and KITTI, which limits the evaluation of richer geometric priors encoded in DA-L [68]. Nevertheless, our prompting strategy effectively provides strong scale priors that narrow the performance gap between backbones of different capacities, enabling consistent metric prediction across MDE sizes.

Beyond the Depth Anything family, we also conduct experiments using ViT-based MoGe-2-Base [62] as the MDE backbone to verify the generality of our scale-prompting design. Tab. 3 shows that SAPE can be effectively applied to MoGe-2 without architectural modification, consistently improving over the original MoGe-2 across all evaluated datasets and depth patterns. This demonstrates that our scale-prompting design is not specific to Depth Anything.

Robustness to Depth Sparsity. To evaluate robustness under varying depth sparsity levels, we conduct experiments on VIO-based sparse points (VOID) and LiDAR scans (KITTI) following standard sparsity protocols. As shown in Tab. 4, our method exhibits remarkable stability across diverse density levels. Notably, despite not being explicitly trained on sparse LiDAR data, both our method and TestPromptDC achieve competitive results on 4-line LiDAR scans compared to OMNI-DC. This underscores the benefits derived from the strong inherent geometric priors of MDE. Furthermore, while TestPromptDC suffers from significant performance degradation in extremely sparse VIO settings (e.g., VOID-150) due to unstable test-time adaptation, our approach maintains high precision without performance collapse, further confirming its superior robustness in depth completion tasks.

Robustness to Depth Range. To evaluate robustness under varying depth range, we conduct experiments on the indoor NYU-Depth V2 and outdoor KITTI datasets using different valid depth ranges. As shown in Fig. 4, our method Any2Full exhibits remarkable consistency across all depth intervals. In

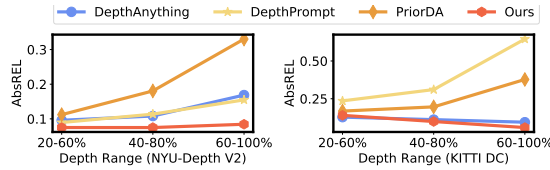


Fig. 4: Robustness comparison on varying depth range. Evaluation is conducted using AbsREL (↓) on NYU-Depth V2 and KITTI DC.

Table 5: Ablation study of Any2Full based on DA-B. **SP**: scale prompt fusion, **LE**: local enrichment, **GP**: global propagation. Metric is AbsREL (↓) evaluated on IBims-1 and KITTI DC across different depth patterns and sparsity levels.

Model	SP	LE	GP	IBims-1			KITTI DC			IBims-1			KITTI DC		
				Hole	Range	Sparse	Hole	Range	Sparse	10%	1%	0.1%	64L	16L	4L
DA-B				0.047	0.048	0.047	0.079	0.087	0.079	0.044	0.044	0.044	0.086	0.087	0.089
Any2Full	✓			0.010	0.068	0.022	0.017	<u>0.089</u>	<u>0.039</u>	0.011	0.020	<u>0.036</u>	0.021	0.032	0.062
Any2Full	✓	✓		<u>0.009</u>	<u>0.066</u>	<u>0.020</u>	<u>0.015</u>	0.092	0.044	<u>0.011</u>	<u>0.017</u>	0.039	<u>0.021</u>	<u>0.032</u>	0.065
Any2Full	✓	✓	✓	0.008	0.031	0.013	0.011	0.046	0.021	0.010	0.012	0.018	0.017	0.022	0.034

contrast, existing MDE-based feed-forward methods, such as DepthPrompt and PriorDA, suffer from significant performance degradation when valid depths are confined to narrow ranges (e.g., the 60-100% interval on KITTI). This highlights a fundamental challenge in scale extrapolation when global range information is limited. By leveraging MDE-guided geometric reasoning to propagate local scale cues, Any2Full effectively extrapolates globally consistent scales to unmeasured regions, ensuring robust completion even under constrained depth observation.

4.3 Ablations and Analysis

To assess the contribution of each module, we conduct ablation studies in Tab. 5 by progressively removing Local Enrichment (LE) and Global Propagation (GP) modules. Comparing Any2Full w/o LE&GP with the DA-B baseline, we observe that the scale prompt (SP) fusion effectively provides necessary scale cues, enabling the MDE backbone to adapt for accurate metric prediction across domains. Furthermore, incorporating the LE and GP modules consistently improves accuracy under varying depth patterns and sparsity levels. Detailed module analyses and additional ablations are provided in the Supp. A.4.

5 Conclusion

In this paper, we present Any2Full, a domain-general and pattern-agnostic depth completion model in one stage. Unlike traditional RGB-D fusion methods, Any2Full reformulates depth completion as a scale-prompting adaptation of monocular depth estimation. It inherits the domain-general geometric priors of a pretrained MDE backbone while flexibly adapting to diverse sparse depth patterns through a scale-aware prompt encoder. Any2Full achieves superior generalization, speed, and accuracy, with practical effectiveness demonstrated in real-world deployment, offering a new perspective on efficient and unified depth completion.

References

1. Arampatzakis, V., Pavlidis, G., Mitianoudis, N., Papamarkos, N.: Monocular depth estimation: A thorough review. *IEEE Transactions on Pattern Analysis and Machine Intelligence* **46**(4), 2396–2414 (2023)
2. Armeni, I., Sax, S., Zamir, A.R., Savarese, S.: Joint 2d-3d-semantic data for indoor scene understanding. *arXiv preprint arXiv:1702.01105* (2017)
3. Bhat, S.F., Birkel, R., Wofk, D., Wonka, P., Müller, M.: Zoedepth: Zero-shot transfer by combining relative and metric depth. *arXiv preprint arXiv:2302.12288* (2023)
4. Bochkovskii, A., Delaunoy, A., Germain, H., Santos, M., Zhou, Y., Richter, S.R., Koltun, V.: Depth pro: Sharp monocular metric depth in less than a second. *arXiv preprint arXiv:2410.02073* (2024)
5. Brown, T., Mann, B., Ryder, N., Subbiah, M., Kaplan, J.D., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., et al.: Language models are few-shot learners. *Advances in neural information processing systems* **33**, 1877–1901 (2020)
6. Cabon, Y., Murray, N., Humenberger, M.: Virtual kitti 2. *arXiv preprint arXiv:2001.10773* (2020)
7. Chen, P.Y., Liu, A.H., Liu, Y.C., Wang, Y.C.F.: Towards scene understanding: Unsupervised monocular depth estimation with semantic-aware representation. In: *Proceedings of the IEEE/CVF Conference on computer vision and pattern recognition*. pp. 2624–2632 (2019)
8. Cheng, J., Liu, L., Xu, G., Wang, X., Zhang, Z., Deng, Y., Zang, J., Chen, Y., Cai, Z., Yang, X.: Monster: Marry monodepth to stereo unleashes power. In: *Proceedings of the Computer Vision and Pattern Recognition Conference*. pp. 6273–6282 (2025)
9. Cheng, X., Wang, P., Guan, C., Yang, R.: Cspn++: Learning context and resource aware convolutional spatial propagation networks for depth completion. In: *Proceedings of the AAAI conference on artificial intelligence*. vol. 34, pp. 10615–10622 (2020)
10. Cheng, X., Wang, P., Yang, R.: Depth estimation via affinity learned with convolutional spatial propagation network. In: *Proceedings of the European conference on computer vision (ECCV)*. pp. 103–119 (2018)
11. Cheng, X., Wang, P., Yang, R.: Learning depth with convolutional spatial propagation network. *IEEE transactions on pattern analysis and machine intelligence* **42**(10), 2361–2379 (2019)
12. Devlin, J., Chang, M.W., Lee, K., Toutanova, K.: Bert: Pre-training of deep bidirectional transformers for language understanding. In: *Proceedings of the 2019 conference of the North American chapter of the association for computational linguistics: human language technologies, volume 1 (long and short papers)*. pp. 4171–4186 (2019)
13. Dosovitskiy, A.: An image is worth 16x16 words: Transformers for image recognition at scale. *arXiv preprint arXiv:2010.11929* (2020)
14. Du, G., Wang, K., Lian, S., Zhao, K.: Vision-based robotic grasping from object localization, object pose estimation to grasp estimation for parallel grippers: a review. *Artificial Intelligence Review* **54**(3), 1677–1734 (2021)
15. Eigen, D., Puhrsch, C., Fergus, R.: Depth map prediction from a single image using a multi-scale deep network. *Advances in neural information processing systems* **27** (2014)
16. Foix, S., Alenya, G., Torras, C.: Lock-in time-of-flight (tof) cameras: A survey. *IEEE Sensors Journal* **11**(9), 1917–1926 (2011)

17. Geiger, A., Lenz, P., Urtasun, R.: Are we ready for autonomous driving? the kitti vision benchmark suite. In: 2012 IEEE conference on computer vision and pattern recognition. pp. 3354–3361. IEEE (2012)
18. Gui, M., Schusterbauer, J., Prestel, U., Ma, P., Kotovenko, D., Grebenkova, O., Baumann, S.A., Hu, V.T., Ommer, B.: Depthfm: Fast generative monocular depth estimation with flow matching. In: Proceedings of the AAAI Conference on Artificial Intelligence. pp. 3203–3211 (2025)
19. Hyoseok, L., Kim, K.S., Byung-Ki, K., Oh, T.H.: Zero-shot depth completion via test-time alignment with affine-invariant depth prior. In: Proceedings of the AAAI Conference on Artificial Intelligence. vol. 39, pp. 3877–3885 (2025)
20. Jeong, C., Bae, I., Park, J.H., Jeon, H.G.: Test-time prompt tuning for zero-shot depth completion. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 9443–9454 (2025)
21. Jeong, C., Bae, I., Park, J.H., Jeon, H.G.: Test-time prompt tuning for zero-shot depth completion. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 9443–9454 (2025)
22. Jia, M., Tang, L., Chen, B.C., Cardie, C., Belongie, S., Hariharan, B., Lim, S.N.: Visual prompt tuning. In: European conference on computer vision. pp. 709–727. Springer (2022)
23. Jun, J., Lee, J.H., Kim, C.S.: Masked spatial propagation network for sparsity-adaptive depth refinement. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 19768–19778 (2024)
24. Ke, B., Obukhov, A., Huang, S., Metzger, N., Daut, R.C., Schindler, K.: Repurposing diffusion-based image generators for monocular depth estimation. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 9492–9502 (2024)
25. Kingma, D.P.: Adam: A method for stochastic optimization. arXiv preprint arXiv:1412.6980 (2014)
26. Koch, T., Liebel, L., Fraundorfer, F., Korner, M.: Evaluation of cnn-based single-image depth estimation methods. In: Proceedings of the European Conference on Computer Vision (ECCV) Workshops. pp. 0–0 (2018)
27. Koch, T., Liebel, L., Fraundorfer, F., Korner, M.: Evaluation of cnn-based single-image depth estimation methods. In: Proceedings of the European Conference on Computer Vision (ECCV) Workshops. pp. 0–0 (2018)
28. Li, X.L., Liang, P.: Prefix-tuning: Optimizing continuous prompts for generation. arXiv preprint arXiv:2101.00190 (2021)
29. Lin, H., Peng, S., Chen, J., Peng, S., Sun, J., Liu, M., Bao, H., Feng, J., Zhou, X., Kang, B.: Prompting depth anything for 4k resolution accurate metric depth estimation. In: Proceedings of the Computer Vision and Pattern Recognition Conference. pp. 17070–17080 (2025)
30. Lin, Y., Cheng, T., Zhong, Q., Zhou, W., Yang, H.: Dynamic spatial propagation network for depth completion. In: Proceedings of the aaai conference on artificial intelligence. vol. 36, pp. 1638–1646 (2022)
31. Liu, X., Shao, X., Wang, B., Li, Y., Wang, S.: Graphcspn: Geometry-aware depth completion via dynamic gens. In: European Conference on Computer Vision. pp. 90–107. Springer (2022)
32. Loshchilov, I., Hutter, F.: Sgdr: Stochastic gradient descent with warm restarts. arXiv preprint arXiv:1608.03983 (2016)
33. Ma, B., Yang, J., Di, D., Zhang, X., Cui, J., Li, H., Xie, Y., Chen, W.: Metricanything: Scaling metric depth pretraining with noisy heterogeneous sources. arXiv preprint arXiv:2601.22054 (2026)

34. Ma, F., Karaman, S.: Sparse-to-dense: Depth prediction from sparse depth samples and a single image. In: 2018 IEEE international conference on robotics and automation (ICRA). pp. 4796–4803. IEEE (2018)
35. Mahler, J., Matl, M., Satish, V., Danielczuk, M., DeRose, B., McKinley, S., Goldberg, K.: Learning ambidextrous robot grasping policies. *Science Robotics* **4**(26), eaau4984 (2019)
36. Maier, D., Hornung, A., Bennewitz, M.: Real-time navigation in 3d environments based on depth camera data. In: 2012 12th IEEE-RAS International Conference on Humanoid Robots (Humanoids 2012). pp. 692–697. IEEE (2012)
37. Miao, R., Liu, W., Chen, M., Gong, Z., Xu, W., Hu, C., Zhou, S.: Ocdepth: A depth-aware method for 3d semantic scene completion. *arXiv preprint arXiv:2302.13540* (2023)
38. Moon, J., Bello, J.L.G., Kwon, B., Kim, M.: From-ground-to-objects: Coarse-to-fine self-supervised monocular depth estimation of dynamic objects with ground contact prior. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 10519–10529 (2024)
39. Park, J.H., Jeong, C., Lee, J., Jeon, H.G.: Depth prompting for sensor-agnostic depth estimation. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 9859–9869 (2024)
40. Park, J., Joo, K., Hu, Z., Liu, C.K., So Kweon, I.: Non-local spatial propagation network for depth completion. In: European conference on computer vision. pp. 120–136. Springer (2020)
41. Perez, E., Strub, F., De Vries, H., Dumoulin, V., Courville, A.: Film: Visual reasoning with a general conditioning layer. In: Proceedings of the AAAI conference on artificial intelligence. vol. 32 (2018)
42. Piccinelli, L., Sakaridis, C., Yang, Y.H., Segu, M., Li, S., Abbeloos, W., Van Gool, L.: Unidepthv2: Universal monocular metric depth estimation made simpler. *IEEE Transactions on Pattern Analysis and Machine Intelligence* (2025)
43. Piccinelli, L., Yang, Y.H., Sakaridis, C., Segu, M., Li, S., Van Gool, L., Yu, F.: Unidepth: Universal monocular metric depth estimation. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 10106–10116 (2024)
44. Ranftl, R., Lasinger, K., Hafner, D., Schindler, K., Koltun, V.: Towards robust monocular depth estimation: Mixing datasets for zero-shot cross-dataset transfer. *IEEE transactions on pattern analysis and machine intelligence* **44**(3), 1623–1637 (2020)
45. Rho, K., Ha, J., Kim, Y.: Guideformer: Transformers for image guided depth completion. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 6250–6259 (2022)
46. Roberts, M., Ramapuram, J., Ranjan, A., Kumar, A., Bautista, M.A., Paczan, N., Webb, R., Susskind, J.M.: Hypersim: A photorealistic synthetic dataset for holistic indoor scene understanding. In: Proceedings of the IEEE/CVF international conference on computer vision. pp. 10912–10922 (2021)
47. Schonberger, J.L., Frahm, J.M.: Structure-from-motion revisited. In: Proceedings of the IEEE conference on computer vision and pattern recognition. pp. 4104–4113 (2016)
48. Schops, T., Schonberger, J.L., Galliani, S., Sattler, T., Schindler, K., Pollefeys, M., Geiger, A.: A multi-view stereo benchmark with high-resolution images and multi-camera videos. In: Proceedings of the IEEE conference on computer vision and pattern recognition. pp. 3260–3269 (2017)

49. Silberman, N., Hoiem, D., Kohli, P., Fergus, R.: Indoor segmentation and support inference from rgbd images. In: European conference on computer vision. pp. 746–760. Springer (2012)
50. Tan, P.X., Hoang, D.C., Nguyen, A.N., Nguyen, V.T., Vu, V.D., Nguyen, T.U., Hoang, N.A., Phan, K.T., Tran, D.T., Vu, D.Q., et al.: Attention-based grasp detection with monocular depth estimation. *IEEE Access* **12**, 65041–65057 (2024)
51. Tang, J., Tian, F.P., An, B., Li, J., Tan, P.: Bilateral propagation network for depth completion. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 9763–9772 (2024)
52. Tang, J., Tian, F.P., Feng, W., Li, J., Tan, P.: Learning guided convolutional network for depth completion. *IEEE Transactions on Image Processing* **30**, 1116–1129 (2020)
53. Tang, Y., Zhao, C., Wang, J., Zhang, C., Sun, Q., Zheng, W.X., Du, W., Qian, F., Kurths, J.: Perception and navigation in autonomous systems in the era of learning: A survey. *IEEE Transactions on Neural Networks and Learning Systems* **34**(12), 9604–9624 (2022)
54. Uhrig, J., Schneider, N., Schneider, L., Franke, U., Brox, T., Geiger, A.: Sparsity invariant cnns. In: 2017 international conference on 3D Vision (3DV). pp. 11–20. IEEE (2017)
55. Vasiljevic, I., Kolkin, N., Zhang, S., Luo, R., Wang, H., Dai, F.Z., Daniele, A.F., Mostajabi, M., Basart, S., Walter, M.R., et al.: Diode: A dense indoor and outdoor depth dataset. *arXiv preprint arXiv:1908.00463* (2019)
56. Viola, M., Qu, K., Metzger, N., Ke, B., Becker, A., Schindler, K., Obukhov, A.: Marigold-dc: Zero-shot monocular depth completion with guided diffusion. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 5359–5370 (2025)
57. Wang, H., Xiao, A., Zhang, X., Yang, M., Lu, S.: Pacgdc: Label-efficient generalizable depth completion with projection ambiguity and consistency. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 7709–7720 (2025)
58. Wang, H., Yang, M., Zheng, N.: G2-monodepth: A general framework of generalized depth inference from monocular rgb+ x data. *IEEE Transactions on Pattern Analysis and Machine Intelligence* **46**(5), 3753–3771 (2023)
59. Wang, H., Che, Z., Yang, Y., Wang, M., Xu, Z., Qiao, X., Qi, M., Feng, F., Tang, J.: Rdgc-gan: Rgb-depth fusion cyclegan for indoor depth completion. *IEEE Transactions on Pattern Analysis and Machine Intelligence* **46**(11), 7088–7101 (2024)
60. Wang, H., Wang, M., Che, Z., Xu, Z., Qiao, X., Qi, M., Feng, F., Tang, J.: Rgb-depth fusion gan for indoor depth completion. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 6209–6218 (2022)
61. Wang, R., Xu, S., Dai, C., Xiang, J., Deng, Y., Tong, X., Yang, J.: Moge: Unlocking accurate monocular geometry estimation for open-domain images with optimal training supervision. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 5261–5271 (2025)
62. Wang, R., Xu, S., Dong, Y., Deng, Y., Xiang, J., Lv, Z., Sun, G., Tong, X., Yang, J.: Moge-2: Accurate monocular geometry with metric scale and sharp details. *Advances in Neural Information Processing Systems* **38**, 35928–35959 (2026)
63. Wang, W., Zhu, D., Wang, X., Hu, Y., Qiu, Y., Wang, C., Hu, Y., Kapoor, A., Scherer, S.: Tartanair: A dataset to push the limits of visual slam. In: 2020 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS). pp. 4909–4916. IEEE (2020)

64. Wang, Y., Li, J., Hong, C., Li, R., Sun, L., Song, X., Wang, Z., Cao, Z., Lin, G.: Tacodepth: Towards efficient radar-camera depth estimation with one-stage fusion. In: Proceedings of the Computer Vision and Pattern Recognition Conference. pp. 10523–10533 (2025)
65. Wang, Z., Chen, S., Yang, L., Wang, J., Zhang, Z., Zhao, H., Zhao, Z.: Depth anything with any prior. arXiv preprint arXiv:2505.10565 (2025)
66. Wong, A., Fei, X., Tsuei, S., Soatto, S.: Unsupervised depth completion from visual inertial odometry. *IEEE Robotics and Automation Letters* **5**(2), 1899–1906 (2020)
67. Yang, L., Kang, B., Huang, Z., Xu, X., Feng, J., Zhao, H.: Depth anything: Unleashing the power of large-scale unlabeled data. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 10371–10381 (2024)
68. Yang, L., Kang, B., Huang, Z., Zhao, Z., Xu, X., Feng, J., Zhao, H.: Depth anything v2. *Advances in Neural Information Processing Systems* **37**, 21875–21911 (2024)
69. Yin, W., Zhang, C., Chen, H., Cai, Z., Yu, G., Wang, K., Chen, X., Shen, C.: Metric3d: Towards zero-shot metric 3d prediction from a single image. In: Proceedings of the IEEE/CVF international conference on computer vision. pp. 9043–9053 (2023)
70. Yu, Z., Zhang, R., Qiu, L., Cao, S.Y., Qiu, K., He, Y., Zhu, S., Dong, Z., Shen, H.L., et al.: Large depth completion model from sparse observations. In: The Fourteenth International Conference on Learning Representations (2026)
71. Zeng, Z., Wang, D., Yang, F., Park, H., Soatto, S., Lao, D., Wong, A.: Wordepth: Variational language prior for monocular depth estimation. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 9708–9719 (2024)
72. Zhang, Y., Guo, X., Poggi, M., Zhu, Z., Huang, G., Mattoccia, S.: Completion-former: Depth completion with convolutions and vision transformers. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 18527–18536 (2023)
73. Zhu, R., Wang, C., Song, Z., Liu, L., Zhang, T., Zhang, Y.: Scaleddepth: Decomposing metric depth estimation into scale prediction and relative depth estimation. arXiv preprint arXiv:2407.08187 (2024)
74. Zuo, Y., Yang, W., Ma, Z., Deng, J.: Omni-dc: Highly robust depth completion with multiresolution depth integration. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 9287–9297 (2025)