

GenHOI: Generalized Hand-Object Pose Estimation with Occlusion Awareness

Hui Yang^{1,2}, Wei Sun^{1*}, Jian Liu^{1*}, Jian Xiao¹, Tao Xie¹,
Hossein Rahmani³, Ajmal Mian⁴, Nicu Sebe⁵, and Gim Hee Lee²

¹ NERC-RVC, Hunan University, China ² National University of Singapore, Singapore

³ Lancaster University, UK ⁴ The University of Western Australia, Australia

⁵ University of Trento, Italy

Abstract. Generalized 3D hand-object pose estimation from a single RGB image remains challenging due to the large variations in object appearances and interaction patterns, especially under heavy occlusion. We propose GenHOI, a framework for generalized hand-object pose estimation with occlusion awareness. GenHOI integrates hierarchical semantic knowledge with hand priors to enhance model generalization under challenging occlusion conditions. Specifically, we introduce a hierarchical semantic prompt that encodes object states, hand configurations, and interaction patterns via textual descriptions. This enables the model to learn abstract high-level representations of hand-object interactions for generalization to unseen objects and novel interactions while compensating for missing or ambiguous visual cues. To enable robust occlusion reasoning, we adopt a multi-modal masked modeling strategy over RGB images, predicted point clouds, and textual descriptions. Moreover, we leverage hand priors as stable spatial references to extract implicit interaction constraints. This allows reliable pose inference even under significant variations in object shapes and interaction patterns. Extensive experiments on the challenging DexYCB and HO3Dv2 benchmarks demonstrate that our method achieves state-of-the-art performance in hand-object pose estimation.

Keywords: Pose Estimation · Hand-Object Interaction

1 Introduction

Accurate 3D hand-object pose recovery is essential for applications in augmented reality [5, 41], human-robot interaction [2, 33, 35], and robotic manipulation [22, 23, 25, 44]. However, real-world hand-object interactions involve a wide variety of objects and interaction scenarios, which hinders the model ability to generalize to unseen objects and interaction patterns. Additionally, frequent occlusions during interactions obscure key visual cues further exacerbating the challenge of generalization in pose estimation.

* Corresponding authors.

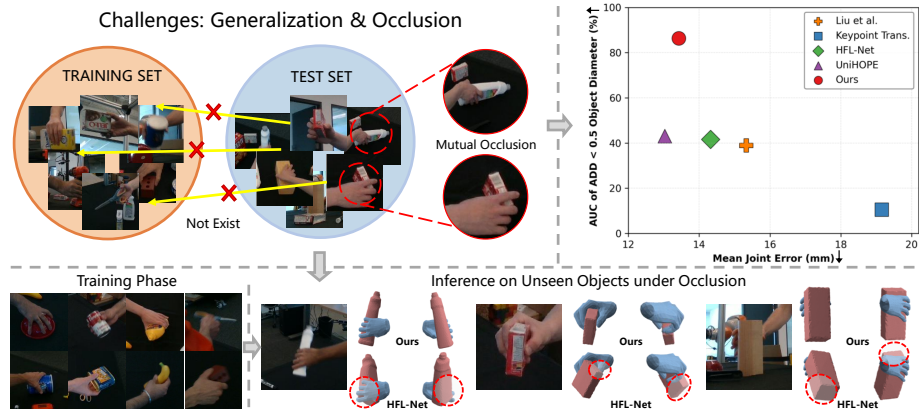


Fig. 1: Illustration of task challenges and key performance comparisons. Hand-object pose estimation (*left*) faces the challenge of generalizing from training set to unseen test set. Additionally, severe hand-object occlusion makes this task more complicated. Our method shows stronger generalization (*bottom*: qualitative comparison with HFL-Net [21] on unseen objects in front and back views, and red dotted circles marking focus areas) and achieves the best performance (*right*) on the challenging generalized DexYCB S3 split [3]. These results confirm effectiveness in generalized settings and under occlusion.

Most existing hand-object pose estimation methods either directly regress poses from visual features [18, 32, 43] or estimate them through explicit keypoint correspondences [19, 21, 31]. The former relies only on appearance cues to capture surface patterns rather than intrinsic object properties or interaction intent. This limits generalization to unseen categories and diverse interaction scenarios. The latter depends on accurate keypoint detection, which makes it highly sensitive to occlusions and missing visual evidence. In general, as shown in Fig. 1, the bottleneck limiting the application of hand-object pose estimation methods lies in their generalization and robustness to occlusion.

To address these limitations, we introduce hierarchical semantic information and interaction priors. Hierarchical semantics provide multi-level descriptions of object properties, hand states, and hand-object interaction patterns. This abstract knowledge supports reasoning about unseen objects and novel interactions, and can supplement missing visual cues. Moreover, the hand provides stable structural priors such as joint positions and rotations during interaction. These internal constraints act as reliable spatial references when object shapes vary widely, which further improves generalization. Consequently, fully leveraging semantic cues and interaction priors is essential to overcome the remaining challenges and to improve the generalization of hand-object pose estimation under occlusion.

In this paper, we propose a framework that integrates hierarchical semantic knowledge with hand priors to improve generalization in hand-object pose

estimation, especially under occlusion. Specifically, we introduce a hierarchical semantic prompt that encodes object states, hand configurations, and interaction patterns through textual descriptions. This prompt guides the model toward abstract representations that generalize to unseen objects and novel interactions. Additionally, we introduce a multi-modal masked modeling strategy that leverages RGB images, predicted point clouds, and language descriptions. By randomly masking and reconstructing regions within these modalities, the model learns to perceive occlusions and infer missing visual cues. This occlusion-aware learning process enables it to maintain reliable pose estimation even when visual information is partially missing or corrupted. Furthermore, we incorporate hand priors by using joint positions and rotational parameters as stable references that capture implicit spatial constraints during interaction. These priors improve robustness across large variations in object shapes and interaction patterns, and they support reliable pose estimation under severe occlusion.

Our contributions are summarized as follows:

- We propose a generalized hand-object pose estimation framework tailored to unseen objects and novel interaction patterns under severe occlusion.
- We introduce hierarchical semantic prompts that encode object states, hand configurations, and interaction patterns through textual descriptions. This supports pose estimation under occlusion with stronger generalization.
- We design a multi-modal masked modeling strategy over RGB images, predicted point clouds, and textual descriptions to learn occlusion-aware representations and infer missing visual cues.
- We leverage hand priors as stable spatial references to extract implicit constraints. This helps the model infer object poses across variations in object shapes and interaction patterns, which improves generalization.

2 Related Work

2.1 3D Hand-Object Pose Estimation

Existing hand-object pose estimation approaches can be broadly divided into direct regression and keypoint-based methods. Direct regression methods directly predict the 3D poses of the hand and object from image features. Chen *et al.* [7] proposed a unified framework that jointly learns image features and parametric hand-object representations to infer poses. Chen *et al.* [6] further introduced kinematic constraints to enforce physically plausible hand motion. Qi *et al.* [29] designed a fused image-depth representation to jointly regress hand-object configurations, while Zhang *et al.* [42] incorporated ray-based feature aggregation to better model local hand-object interactions. Keypoint-based approaches estimate the 3D hand-object pose by leveraging spatial relations among detected keypoints. Doosti *et al.* [11] jointly predicted hand and object keypoints and used MANO [30] and PnP to recover 3D poses. Liu *et al.* [26] introduced a semi-supervised framework that utilizes spatiotemporal constraints to provide pseudo-label supervision. Hampali *et al.* [13] integrated cross-attention transformers to

explicitly encode joint hand-object geometry, and Lin *et al.* [21] further enforced dual-stream feature disentanglement for more stable mutual pose refinement.

Despite the progress of both paradigms, most existing methods rely primarily on visual cues or fixed keypoint representations and lack semantic and structural priors. As a result, they struggle in heavy-occlusion scenarios and exhibit limited generalization to unseen objects and novel interaction patterns.

2.2 Textual Semantics for Hand-Object Interaction

Recent studies [27, 40] indicate that purely visual cues often struggle under severe occlusion and large intra-class variation, where textual semantic can provide effective complementary guidance. This has motivated the integration of textual semantics into interaction modeling. Zhang *et al.* [45] used large language models to produce contextual interaction descriptions and combine them with image generators to synthesize hand-object background images. Following this direction, Cha *et al.* [1] introduced a text-guided 3D interaction generation pipeline to produce plausible hand-object configurations. Huang *et al.* [17] further build a generative language-to-3D framework that enables interpretable and expressive interaction synthesis. In parallel, textual supervision has been adopted in object pose estimation [24, 37, 39]. Corsetti *et al.* [8] leveraged vision-language priors for open-vocabulary segmentation and 6D pose estimation, while Lin *et al.* [20] leveraged pretrained vision-language models to incorporate semantic knowledge from both image and text modalities for enhanced category-level understanding.

However, these methods mainly rely on textual cues that describe the global scene or coarse object categories, which provide only high-level context. In contrast, our approach tackles the more challenging hand-object pose estimation task, which demands fine-grained spatial reasoning under severe occlusion. We introduce a hierarchical textual representation that encodes the object, the hand, and their interaction to provide structured multi-level cues for detailed understanding. These hierarchical semantics are integrated with visual and geometric modalities through a multimodal masked modeling strategy, which reconstructs masked signals under random occlusion and improves occlusion awareness and cross-modal reasoning.

3 Method

Fig. 2 shows our proposed GenHOI, a novel framework that integrates hierarchical semantic prompts and hand priors to enhance both generalization and robustness under occlusion.

3.1 Hierarchical Textual Semantic Embedding

We construct a hierarchical prompt template that encodes multi-level semantics of the hand, object, and their interaction. Leveraging this template with the input image, we generate a unified textual prompt that captures both fine-grained visual details and high-level semantic relationships.

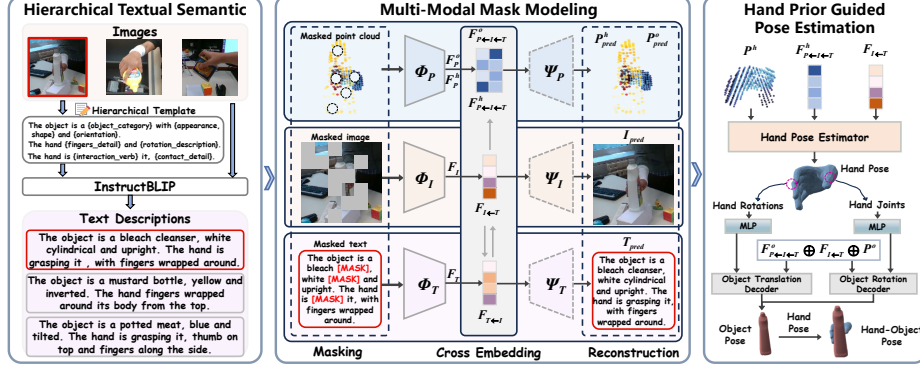


Fig. 2: GenHOI framework, which consists of three key components: (1) **Hierarchical Textual Semantics**. Given RGB images and hierarchical templates with the object, hand, and interaction levels, we employ InstructBLIP to generate hierarchical textual descriptions. (2) **Multi-Modal Mask Modeling**. Given the RGB image, textual description, and corresponding hand-object point cloud, we first apply a modality-specific masking strategy to the inputs. The masked inputs are processed by the cross-modal embedding module, which fuses visual, geometric, and textual features to produce representations used for reconstruction and pose estimation. The dashed boxes indicate that masking and reconstruction are used only during training. (3) **Hand Prior Guided Pose Estimation**. The cross-modal features learned in the previous stage are aggregated to a robust representation. Using these fused features, we first estimate hand pose parameters, which then serve as reliable priors for object pose reasoning.

Hierarchical Prompt Template. We define three semantic components to construct the hierarchical prompt template: object-level semantics s_o , hand-level semantics s_h , and interaction-level semantics s_i .

The object-level semantics s_o describe object properties such as category, appearance, and orientation, formulated as “the object is a {object_category} with a {shape, appearance}, and {orientation_description}.” The hand-level semantics s_h characterize hand articulation and grasp configuration, expressed as “the hand {fingers_detail} and {rotation_description}.” The interaction-level semantics s_i capture the relational context between the hand and the object, following the structure “the hand is {interaction_verb} it, {contact_detail}.” Finally, these three components are integrated into a unified hierarchical prompt $\mathcal{T}$. This process enables the model to capture a comprehensive multi-level semantic context by linking cues from the object, hand, and their interaction. The final hierarchical template is shown in Fig. 2.

Hierarchical Textual Semantic Generation. Given an input image $I \in \mathbb{R}^{3 \times H \times W}$ and a hierarchical template $\mathcal{T}$, we leverage a vision-language model such as InstructBLIP [9] to generate a textual description T . Formally, the generation process is defined as:

$$T = \mathcal{G}(I, \mathcal{T}), \quad (1)$$

where $\mathcal{G}(\cdot)$ denotes the generative vision-language model. This model fills the template slots with image-specific content to produce a textual description that aligns the visual observations with the hierarchical semantics encoded in $\mathcal{T}$. As shown in Fig. 2, the model generates the hierarchical prompt $\mathcal{T}$ based on the visual content. To improve robustness against potential noise or bias in automatically generated descriptions, textual perturbations are introduced during training to reduce dependence on perfectly consistent inputs. More details are provided in the supplementary material.

This unified hierarchical textual description provides rich semantic information that encapsulates both local and global contextual cues of the hand-object interaction. It is then fused with visual and geometric representations to enable a more robust and generalizable pose estimation under occlusion.

3.2 Multi-Modal Masked Modeling

Visual cues often become incomplete under severe occlusion due to hand object interaction, while point cloud geometry and textual semantics can provide complementary information. To enhance occlusion awareness and improve cross-modal consistency, we propose a multi-modal masked modeling strategy that jointly masks and reconstructs visual, geometric, and textual modalities. During training, each image is paired with a point cloud sampled from ground-truth hand and object mesh. During inference, we dynamically sample a point cloud from a voxel grid guided by image features. The sampling follows [29, 38] and provides a geometric representation for each image without requiring ground-truth meshes. Masking and reconstruction are applied only during training to help the model infer missing information and align multi-modal representations. Moreover, the model directly utilizes the learned cross-modal fusion features from image, point cloud, and text, and these integrated representations are then used for subsequent hand-object pose estimation.

Modality-specific Masking Strategy. We apply three modality-specific masks:

- 1) Image masking corrupts selected patches in the hand-object region and the background.
- 2) Point cloud masking drops random local patches for both hand and object geometry.
- 3) Text masking replaces a subset of tokens in the hierarchical description with [MASK].

Image Masking: Given an input image I , the masked image $\tilde{I}$ is defined as:

$$\begin{aligned} M_I &= M_{ho}(\lfloor \alpha N_I \rfloor) \cup M_{bg}(N_I - \lfloor \alpha N_I \rfloor), \\ \tilde{I} &= (1 - M_I) \odot I + M_I \odot \mathcal{N}(0, \sigma^2), \end{aligned} \quad (2)$$

where $\odot$ denotes element-wise multiplication, and $M_I \in \{0, 1\}^{H \times W}$ is the mask for the image. Here, M_{ho} and M_{bg} are the sets of masked patches sampled from the hand-object region and background region, respectively. N_I is the total number of masked patches and $\alpha \in [0, 1]$ controls the proportion of patches masked

within the hand-object region. $\mathcal{N}(0, \sigma^2)$ denotes Gaussian noise. Our target-centric masking balances patches from the hand-object region and background, which preserves interaction-focused cues and global contextual information.

Point Cloud Masking: Given a hand point cloud $P^h \in \mathbb{R}^{3 \times N_h}$, we first partition it into K local patches $\bigcup_{k=1}^K \mathbf{P}_k^h = P^h$. The masked hand point cloud $\tilde{\mathbf{P}}_k^h$ is then defined as:

$$\tilde{\mathbf{P}}_k^h = \begin{cases} \mathbf{0}, & k \in \mathcal{K}^h, \\ \mathbf{P}_k^h, & k \notin \mathcal{K}^h, \end{cases} \quad (3)$$

where $\mathcal{K}^h$ is the set of patch indices randomly selected according to the masking ratio $\beta \in [0, 1]$. The object point cloud P^o is masked in the same manner as hand, yielding the final masked point clouds $\tilde{P}^h$ and $\tilde{P}^o$.

Textual Masking: Given a hierarchical textual description $\mathbf{T} = \{t_l\}_{l=1}^L$ of length L , we randomly select a subset of token positions $\mathcal{L}$ to mask according to a masking ratio $\gamma \in [0, 1]$. The masked sequence $\tilde{\mathbf{T}} = \{\tilde{t}_l\}_{l=1}^L$ is then defined as:

$$\tilde{t}_l = \begin{cases} [\text{MASK}], & l \in \mathcal{L}, \\ t_l, & l \notin \mathcal{L}, \end{cases} \quad (4)$$

where [MASK] denotes the special mask token [10].

Multi-Modal Masked Reconstruction. Given the masked inputs $\tilde{I}$, $\tilde{P}$, and $\tilde{T}$ obtained from the previous stage, the model first extracts modality-specific representations to encode visual, geometric, and textual information. Specifically, the masked image $\tilde{I}$ is encoded by the visual encoder ϕ_I to produce $F_I \in \mathbb{R}^{H_I \times W_I \times D_I}$. The masked point cloud $\tilde{P}$ is encoded by the geometric encoder ϕ_P to produce $F_P \in \mathbb{R}^{N_P \times D_P}$. The masked textual description $\tilde{T}$ is encoded by the textual encoder ϕ_T to produce $F_T \in \mathbb{R}^{L_T \times D_T}$. These representations serve as the foundational features for subsequent cross-modal embedding and reconstruction.

Visual and textual features are first fused through cross-attention since text-guided visual features help produce more accurate point cloud sampling. For visual feature F_I and textual feature F_T , the attention is formulated as:

$$\begin{aligned} \mathbf{Q}_I &= F_I \mathbf{W}_Q^I, \quad \mathbf{K}_T = F_T \mathbf{W}_K^T, \quad \mathbf{V}_T = F_T \mathbf{W}_V^T, \\ F_{I \leftarrow T} &= \text{Softmax}\left(\frac{\mathbf{Q}_I \mathbf{K}_T^\top}{\sqrt{d}}\right) \mathbf{V}_T, \end{aligned} \quad (5)$$

with $F_{T \leftarrow I}$ obtained through the same formulation as above. Reconstruction heads ψ_I and ψ_T predict the recovered image I_{pred} and textual description T_{pred} :

$$I_{pred} = \psi_I(F_{I \leftarrow T}), \quad T_{pred} = \psi_T(F_{T \leftarrow I}). \quad (6)$$

The semantically enhanced visual feature $F_{I \leftarrow T}$ is fused with the hand point cloud feature F_P^h via cross-attention to produce the fused enhanced hand point

cloud feature $F_{P \leftarrow T \leftarrow I}^h$. The enhanced features of the hand point cloud $F_{P \leftarrow I \leftarrow T}^h$ is directly fed into a reconstruction head ψ_P to predict the reconstructed hand point cloud:

$$P_{\text{pred}}^h = \psi_P(F_{P \leftarrow I \leftarrow S}^h). \quad (7)$$

The object point cloud P_{pred}^o is reconstructed in the same manner.

This hierarchical fusion and reconstruction mechanism lets each modality draw on complementary signals from the others. It supports occlusion-aware completion and preserves semantic coherence and structural consistency across visual, geometric, and textual modalities.

3.3 Hand Prior Guided Pose Estimation

To enhance the generalization capability of the model across diverse objects and interaction scenarios, we introduce a hand prior guided hand-object pose estimation module. This module leverages the inherent structural regularities of the hand as reliable prior information to guide object pose reasoning. Compared to diverse and deformable objects, the hand exhibits a relatively fixed motion topology and limited shape variation. The estimated hand pose and geometric information thus provide stable and effective spatial constraints for inferring object pose.

Hand Pose Estimation. We incorporate an implicit geometric by predicting the signed distance field (SDF). Given hand points P_h , we apply Fourier positional encoding $\gamma(\cdot)$ to enrich spatial frequency representation, and fuse it with $F_{I \leftarrow T}$ to regress the SDF_h :

$$SDF_h = \text{MLP}_{\text{SDF}}(\gamma(P_h) \oplus F_{I \leftarrow T} \oplus P_h). \quad (8)$$

The process of estimating object SDF_o follows the same steps as in SDF_h . We estimate the hand pose by integrating multiple cues, including the explicit geometric feature, the semantically enhanced visual feature, and the implicit geometric from the predicted SDF_h . The SDF-guided aggregation encourages features to focus on the hand-object surface.

$$F_{\text{agg}}^h = (F_{P \leftarrow I \leftarrow S}^h \oplus F_{I \leftarrow T}) \cdot \frac{1}{\beta_h} \cdot \sigma_h \left(\frac{SDF_h}{\beta_h} \right), \quad (9)$$

$$(\{\theta_i\}_{i=1}^{16}, \alpha) = \text{MLP}_h \left(\text{Transformer}_h(F_{\text{agg}}^h) \right), \quad (10)$$

where $\sigma_h(\cdot)$ is the sigmoid function, β_h is a learnable scale, $\theta \in \mathbb{R}^{16 \times 3}$ represents joint rotations, and $\alpha \in \mathbb{R}^{10}$ denotes MANO hand shape coefficients. The aggregation object feature F_{agg}^o is obtained in the same manner as for the hand.

Object Pose Estimation. The object rotation and translation are jointly regressed by incorporating hand cues. Specifically, the joint rotations θ are projected and globally aggregated to produce a compact hand motion descriptor, and the reconstructed 3D hand joints $J \in \mathbb{R}^{21 \times 3}$ provide spatial cues of the grasp. These hand features are concatenated with the object feature F_{agg}^o to estimate the full object pose:

$$\begin{aligned} R_{\text{obj}} &= \text{MLP}_R \left(F_{\text{agg}}^o \oplus \text{Pool}(\phi_{\text{rot}}(\theta_{i=1}^{16})) \right), \\ T_{\text{obj}} &= \text{MLP}_T \left(F_{\text{agg}}^o \oplus \text{Pool}(\phi_{\text{pos}}(J)) \right). \end{aligned} \quad (11)$$

The final object pose is represented as $[R_{\text{obj}}, T_{\text{obj}}]$, where $R_{\text{obj}} \in \mathbb{R}^{3 \times 3}$ is a 3D rotation matrix and $T_{\text{obj}} \in \mathbb{R}^3$ is a 3D translation vector.

3.4 Loss Functions

Our overall training objective:

$$\begin{aligned} L_{\text{total}} &= \lambda_1 L_{\text{rec}} + \lambda_2 L_{\text{text}} + \lambda_3 L_{\text{pc}} \\ &\quad + \lambda_4 L_{\text{mano}} + \lambda_5 L_{\text{obj}} + \lambda_6 L_{\text{SDF}} + \lambda_7 L_{\text{others}}. \end{aligned} \quad (12)$$

combines multiple loss terms to jointly optimize multi-modal reconstruction quality and pose estimation accuracy. Specifically, we implement the image reconstruction loss L_{rec} , the point cloud reconstruction loss L_{pc} , the SDF supervision loss L_{sdf} , the hand pose loss L_{mano} , and the object pose loss L_{obj} as L1 losses. These loss terms are supervised by image cues, geometric consistency, implicit geometry, hand pose, and object pose, respectively. In contrast, the text reconstruction loss L_{text} is formulated as a cross-entropy loss to reconstruct masked text tokens. The auxiliary loss terms L_{others} follow Qi *et al* [29] to provide task-specific constraints. Each term is weighted by its corresponding coefficient λ to balance its contribution during optimization.

4 Experiments

Our method is implemented in PyTorch and trained on a single NVIDIA RTX 4090 GPU. We adopt the Adam optimizer with a batch size of 24. The initial learning rate is set to 1e-4 and decayed by a factor of 0.7 every 5 epochs. The model is trained for 60 epochs on both the DexYCB [3] and HO3Dv2 [12] datasets, which is sufficient to achieve satisfactory performance. Further details are in supplementary material.

Datasets. We evaluate our method on two widely used hand-object interaction datasets. **DexYCB** [3] contains 582K images from over 1000 sequences, where 10 subjects manipulate 20 YCB objects [34]. We follow the S0 split for standard evaluation. The split uses sequences from 8 subjects for training and 2 subjects for testing, which keeps the test subjects and interaction patterns unseen during

Table 1: Evaluation of hand and object pose estimation on the DexYCB S3 split [3]. Results are for unseen objects and novel hand-object interactions. ‘avg’ denotes average results over all evaluated objects. ↓ means lower is better, ↑ indicates higher is better. Best results are bolded, and the second best are underlined. HPE: Hand Pose Estimation. HOPE: Hand-Object Pose Estimation.

Method		Object Pose (mAP)				Hand Pose (mm)			
		AUC of ADD-0.5D ↑				MJE ↓ PA-MJE ↓ V-PE ↓ PA-V-PE ↓			
		gelatin_box	bleach_cleanser	wood_block	avg				
HPE	HandOccNet [28]	—	—	—	—	14.58	6.73	14.10	6.49
	MobRecon [4]	—	—	—	—	15.40	7.25	14.24	6.39
	H2ONet [36]	—	—	—	—	15.20	<u>6.35</u>	15.03	6.74
	SimpleHand [46]	—	—	—	—	14.88	6.74	14.21	6.45
HOPE	Liu et al. [26]	<u>26.31</u>	25.07	68.56	38.89	15.43	6.61	14.91	6.40
	Keypoint Trans. [13]	0.00	1.31	32.61	10.47	18.79	7.77	18.35	7.94
	HFL-Net [21]	25.88	32.08	70.16	41.59	14.77	6.64	14.29	6.41
	UniHOPE [32]	26.23	<u>32.32</u>	<u>74.29</u>	<u>43.06</u>	13.31	6.08	12.89	5.87
	GenHOI (ours)	86.70	89.27	92.06	89.34	<u>13.67</u>	6.39	<u>13.32</u>	<u>5.93</u>

training. We also adopt the S3 split to evaluate generalization to unseen objects. The split uses 15 objects for training and holds out the remaining 3 objects for testing. **HO3Dv2** [12] contains 77K images from 68 sequences. The dataset covers 10 subjects interacting with 10 YCB objects. We use the official train/test split, where the test set contains 1 unseen object category to evaluate cross-category generalization, and submit test results to the official evaluation server.

Evaluation Metrics. Hand Pose and Mesh Reconstruction: To assess hand pose accuracy, we report Mean Joint Error (M-JE), Procrustes-Aligned MJE (PA-MJE) [47], and Scale-Translation aligned MJE (ST-MJE) [48]. Since our method directly regresses MANO parameters, we further evaluate hand mesh reconstruction quality using Mean Mesh Error (MME), the area under the curve of the percentage of correct vertices (V-AUC), and F-scores at 5mm and 15mm thresholds (F@5, F@15), together with their Procrustes-aligned versions following H2ONet [36]. **Object Pose:** For object pose estimation, we report Object Center Error (OCE), Mean Corner Error (MCE), Object Mesh Error (OME), and average closest point distance (ADD-S). We further evaluate the average 3D distance (ADD) of the object. Specifically, we report ADD-0.5D, which measures the percentage of predictions whose ADD is within 50% of the object diameter.

Table 2: Evaluation of object pose estimation using ADD-S on DexYCB S3 split [3].

Method	ADD-S (mm) ↓			
	gelatin_box	bleach_cleanser	wood_block	average
HFL-Net [21]	94.74	176.84	72.93	114.83
UniHOPE [32]	109.23	173.38	<u>49.76</u>	110.79
HOISDF [29]	<u>52.31</u>	<u>89.92</u>	109.89	84.04
GenHOI (ours)	13.14	26.91	25.23	21.76

Table 3: Quantitative comparison of hand mesh with across occlusion levels on DexYCB S3 split dataset [3].

Methods	Occlusion (25%-50%)				Occlusion (50%-75%)				Occlusion (75%-100%)			
	J-PE↓	PA-J-PE↓	V-PE↓	PA-V↓	J-PE↓	PA-J-PE↓	V-PE↓	PA-V↓	J-PE↓	PA-J-PE↓	V-PE↓	PA-V↓
HFL-Net [21]	16.33	7.00	15.81	6.77	18.66	7.33	18.11	7.11	28.95	8.80	27.94	8.53
UniHOPE [32]	14.59	6.39	14.13	6.17	16.27	6.51	15.78	6.29	26.42	7.64	25.51	7.40
Ours	<u>14.75</u>	<u>6.49</u>	<u>14.36</u>	<u>6.23</u>	16.18	6.61	<u>15.89</u>	6.25	26.28	7.55	<u>25.64</u>	<u>7.57</u>

Table 4: Quantitative results on the unseen object (“019 pitcher base”) from the HO3Dv2 dataset [12].

Method	HFL-Net [21]	HOISDF [29]	Ours
Metrics in [mm]	ADD-0.5D	ADD-0.5D	ADD-0.5D
019 pitcher base	24.1	88.4	92.7

4.1 Comparison with SOTA Methods

Comparisons on DexYCB. We evaluate our method on the DexYCB S3 split dataset [3], which focuses on unseen objects. Following UniHOPE [32], we evaluate only hand-object interaction frames. This split presents significant challenges due to severe hand-object occlusions and the inclusion of unseen objects, providing a rigorous evaluation of the generalization ability and robustness. Table 1 presents quantitative results for hand and object pose estimation. Our method achieves competitive hand pose accuracy, closely matching state-of-the-art (SOTA) approaches. More importantly, it substantially outperforms existing methods in object pose estimation by achieving an average AUC of 89.34% across all evaluated categories. The relatively smaller improvement in hand pose accuracy compared to object pose stems from the inherent differences in variability. The hand exhibits a consistent structure and appearance across samples, which makes it less influenced by the generalization design of our framework. To enable a clearer comparison under a shared mm-based object metric, we further report the ADD-S results in Table 2. The results show that our method achieves the lowest ADD-S error among all compared methods, reducing the error from 84.04 mm of HOISDF to 21.76 mm. In contrast, objects exhibit much larger shape and texture variation. Our hierarchical textual semantics and hand priors therefore provide stronger semantic and geometric constraints. Following UniHOPE, we evaluate hand pose accuracy across occlusion ratios under the same protocol in Table 3. Our method achieves comparable performance under heavy occlusion despite UniHOPE includes a dedicated hand-occlusion removal module. Fig. 3 shows the corresponding qualitative comparison on unseen objects under the S3 split. The results highlight improved hand-object pose accuracy over HFL-Net and UniHOPE.

Comparisons on HO3Dv2. We evaluate our method on the HO3Dv2 dataset [12], which mainly consists of objects seen during training, along with a single unseen object (“019 pitcher base”). To assess the generalization capability, we

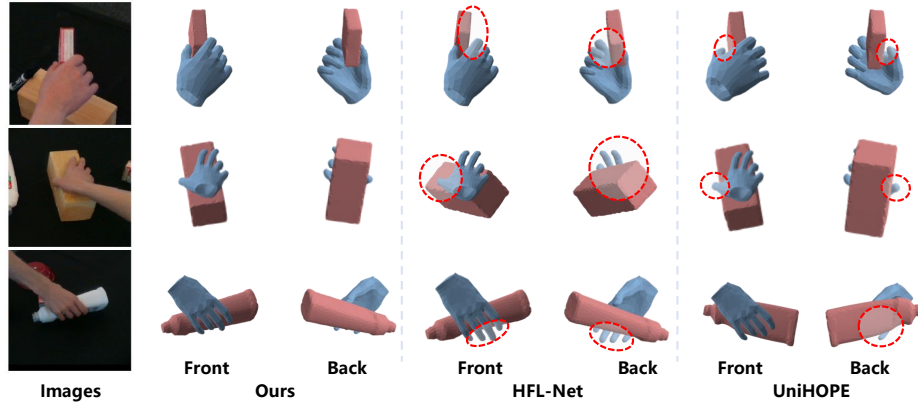


Fig. 3: Qualitative comparison of hand-object pose estimation on unseen objects under the DexYCB S3 split [3]. Front and back indicate the front and rear views, respectively. Red dotted circles highlight regions where other methods produce less accurate pose estimates than our method. This demonstrates the superior generalization and robustness of our method on unseen objects.

Table 5: Comparison with SOTA hand-object pose estimation methods on the HO3Dv2 dataset [12].

Metric in [mm]	MJE ↓	ST-MJE ↓	PA-MJE ↓	OME ↓	ADD-S ↓
Hasson <i>et al.</i> [16]	-	31.8	11.0	-	-
Hasson <i>et al.</i> [14]	-	36.9	11.4	67.0	22.0
Hasson <i>et al.</i> [15]	-	26.8	12.0	80.0	40.0
Liu <i>et al.</i> [26]	-	31.7	10.1	-	-
Hampali <i>et al.</i> [13]	25.5	25.7	10.8	68.0	21.4
HFL-Net [21]	28.9	28.4	8.9	64.3	32.4
HOISDF [29]	<u>23.6</u>	22.8	<u>9.6</u>	<u>48.5</u>	<u>17.8</u>
LCP [19]	-	23.7	8.9	-	-
UniHOPE [32]	-	25.5	<u>9.6</u>	-	-
GenHOI (ours)	21.3	21.1	9.9	35.9	13.4

evaluate our model on the unseen object, as presented in Table 4. Despite not being seen during training, our method achieves a competitive ADD-0.5D score of 92.7% that surpasses existing approaches. This strong generalization to unseen objects benefits from the integration of hierarchical textual semantics and hand priors. Table 5 reports quantitative comparisons on the seen objects. Our method consistently surpasses previous approaches across all metrics in achieving the lowest hand and object pose errors. These results demonstrate that the proposed model accurately captures hand-object interactions and yields robust estimation even under challenging occlusions.

Qualitative comparisons on both unseen and seen objects are shown in Fig. 4. Specifically, the top-left example corresponds to the unseen object, while the re-

Table 6: Ablation study of our proposed components on the DexYCB dataset S3 split full [3]. ‘w/o’ denotes ‘without’.

Ablation	MJE ↓ PA-MJE ↓ V-PE ↓ PA-V-PE ↓ ADD-0.5D ↑				
Hierarchical Textual Semantic					
1 w/o Text Generation Template	16.78	6.73	15.12	6.23	86.51
2 w/o Multi-level Semantic Knowledge	17.04	6.98	15.75	6.41	86.32
Multi-Modal Masked Modeling					
3 w/o Image Masking	15.33	6.52	14.78	6.03	87.65
4 w/o Text Masking	14.98	6.21	14.50	5.88	88.12
5 w/o Point Cloud Masking	15.21	6.47	14.91	6.07	87.90
Hand Priors					
6 w/o Hand Rotation Prior	14.65	6.12	14.32	5.78	85.45
7 w/o Hand Joint Prior	14.88	6.29	14.55	5.91	84.01
GenHOI (Full)	13.42	5.81	12.96	5.63	90.44

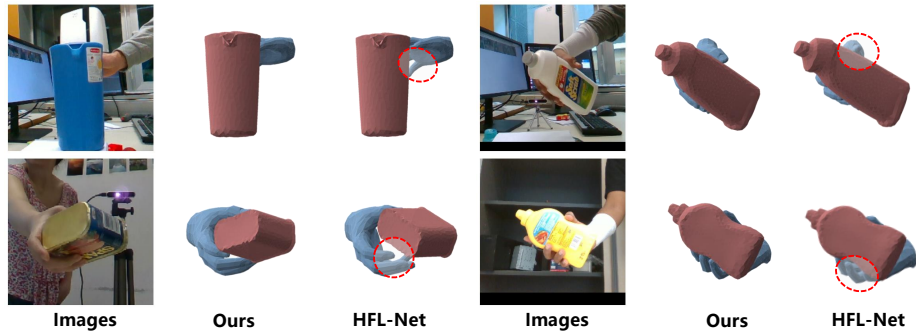


Fig. 4: Qualitative comparison on the HO3Dv2 dataset [12]. The top-left example show results on the unseen object (“019 pitcher base”), while the remaining examples correspond to objects seen during training.

maining examples show results on seen objects. The qualitative results further confirm that our method attains SOTA performance on HO3Dv2 while maintaining robust generalization beyond the training distribution.

4.2 Ablation Studies

We perform the ablations on the DexYCB S3 split full [3], using all frames to evaluate the contribution of each component.

Hierarchical Textual Semantic. Rows (1-2) of Table 6 show apparent performance drops when removing either the text template or hierarchical semantic descriptions. Our motivation for introducing hierarchical semantics is to provide structured level-dependent priors about object and interaction attributes. These cues offer high-level concepts as well as fine-grained part-level hints, which guide

the model to reason beyond purely visual correlation. The model is forced to rely mainly on appearance cues without them, and thus weakening its generalizability. The resulting performance gap confirms that hierarchical semantics serve as a semantic generalization prior, which is crucial when handling novel hand-object interaction scenarios with severe occlusion.

Multi-Modal Masked Modeling. Rows (3-5) of Table 6 indicate that image masking contributes the most to performance, and text and point cloud masking further enhance occlusion awareness. The multi-modal masking strategy enables the model to infer occluded information. The model develops occlusion-aware reasoning and cross-modal feature agreement by learning to reconstruct from incomplete observations, which are crucial for robust performance in complex interaction scenes. The consistent performance degradation when the masks are removed validates that explicitly learning to reason under occlusion is central to improving generalization in real-world settings.

Hand Priors. Rows (6-7) of Table 6 show that the removal of hand priors only slightly affects hand pose estimation but leads to a clear drop in object pose accuracy, which indicates that stable hand priors provide essential spatial constraints for interaction reasoning. This is because hand configuration and articulation inherently encode grasp intent and contact topology, which implicitly restrict the feasible orientation and position of the object in 3D space. Even when the hand pose can still be estimated from appearance cues, the loss of these structured priors weakens the geometric coupling between hand and object. As a result, object pose inference becomes more sensitive to local visual variations. The model relies more on appearance similarity without such constraints. This makes object pose inference harder when shapes and interaction patterns vary, which reduces generalization to unseen objects and novel interaction scenarios.

Additional ablations and analyses on multimodal mask ratios, multimodal fusion strategies, robustness to inaccurate or noisy text descriptions, as well as inference efficiency are provided in the supplementary material.

5 Conclusion

In this work, we presented GenHOI, a novel hand-object pose estimation framework that improves generalization to unseen objects and interaction scenarios. We introduced hierarchical textual semantics to encode object states, hand configurations, and interaction patterns that provide high-level guidance beyond visual appearance. We proposed a multi-modal masked modeling strategy that enables the model to infer occluded regions and maintain cross-modal consistency under partial observations. Furthermore, hand priors serve as stable spatial constraints that anchor interaction geometry to reinforce object pose reasoning even when object shapes vary. Extensive experiments demonstrate that our GenHOI achieves strong generalization and robust performance under severe occlusion. For future work, we may extend the model to egocentric datasets and further validate its generalization.

Acknowledgements

This work is supported by the National Natural Science Foundation of China under Grants 62473141 and U22A2059, National Science and Technology Major Project of China under Grants 2026ZD1610900, China Scholarship Council under Grant 202506130069, the Open Foundation of the State Key Laboratory of Advanced Design and Manufacturing for Vehicle Body, the National Research Foundation (NRF) Singapore, under its NRF-Investigatorship Programme (Award ID. NRF-NRFI09-0008), and the Tier 2 grant MOET2EP20124-0015 from the Singapore Ministry of Education.

References

1. Cha, J., Kim, J., Yoon, J.S., Baek, S.: Text2hoi: Text-guided 3d motion generation for hand-object interaction. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 1577–1585 (2024) [4](#)
2. Chao, Y.W., Paxton, C., Xiang, Y., Yang, W., Sundaralingam, B., Chen, T., Murali, A., Cakmak, M., Fox, D.: Handoversim: A simulation framework and benchmark for human-to-robot object handovers. In: Proceedings of the IEEE International Conference on Robotics and Automation. pp. 6941–6947 (2022) [1](#)
3. Chao, Y.W., Yang, W., Xiang, Y., Molchanov, P., Handa, A., Tremblay, J., Narang, Y.S., Van Wyk, K., Iqbal, U., Birchfield, S., et al.: Dexycb: A benchmark for capturing hand grasping of objects. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 9044–9053 (2021) [2](#), [9](#), [10](#), [11](#), [12](#), [13](#)
4. Chen, X., Liu, Y., Dong, Y., Zhang, X., Ma, C., Xiong, Y., Zhang, Y., Guo, X.: Mobrecon: Mobile-friendly hand mesh reconstruction from monocular image. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 20544–20554 (2022) [10](#)
5. Chen, Y., Wang, Q., Chen, H., Song, X., Tang, H., Tian, M.: An overview of augmented reality technology. In: Journal of Physics: Conference Series. vol. 1237, p. 022082 (2019) [1](#)
6. Chen, Z., Chen, S., Schmid, C., Laptev, I.: Gsdf: Geometry-driven signed distance functions for 3d hand-object reconstruction. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 12890–12900 (2023) [3](#)
7. Chen, Z., Hasson, Y., Schmid, C., Laptev, I.: Alignsdf: Pose-aligned signed distance fields for hand-object reconstruction. In: European Conference on computer vision. pp. 231–248. Springer (2022) [3](#)
8. Corsetti, J., Boscaini, D., Oh, C., Cavallaro, A., Poiesi, F.: Open-vocabulary object 6d pose estimation. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 18071–18080 (2024) [4](#)
9. Dai, W., Li, J., Li, D., Tiong, A., Zhao, J., Wang, W., Li, B., Fung, P.N., Hoi, S.: Instructblip: Towards general-purpose vision-language models with instruction tuning. Advances in Neural Information Processing Systems **36**, 49250–49267 (2023) [5](#)
10. Devlin, J., Chang, M.W., Lee, K., Toutanova, K.: Bert: Pre-training of deep bidirectional transformers for language understanding. In: Proceedings of the conference

- of the North American chapter of the association for computational linguistics: human language technologies. pp. 4171–4186 (2019) [7](#)
11. Doosti, B., Naha, S., Mirbagheri, M., Crandall, D.J.: Hope-net: A graph-based model for hand-object pose estimation. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 6608–6617 (2020) [3](#)
 12. Hampali, S., Rad, M., Oberweger, M., Lepetit, V.: Honnotate: A method for 3d annotation of hand and object poses. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 3196–3206 (2020) [9](#), [10](#), [11](#), [12](#), [13](#)
 13. Hampali, S., Sarkar, S.D., Rad, M., Lepetit, V.: Keypoint transformer: Solving joint identification in challenging hands and object interactions for accurate 3d pose estimation. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 11090–11100 (2022) [3](#), [10](#), [12](#)
 14. Hasson, Y., Tekin, B., Bogo, F., Laptev, I., Pollefeys, M., Schmid, C.: Leveraging photometric consistency over time for sparsely supervised hand-object reconstruction. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 571–580 (2020) [12](#)
 15. Hasson, Y., Varol, G., Schmid, C., Laptev, I.: Towards unconstrained joint hand-object reconstruction from rgb videos. In: International Conference on 3D Vision. pp. 659–668 (2021) [12](#)
 16. Hasson, Y., Varol, G., Tzionas, D., Kalevatykh, I., Black, M.J., Laptev, I., Schmid, C.: Learning joint reconstruction of hands and manipulated objects. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 11807–11816 (2019) [12](#)
 17. Huang, M., Chu, F.J., Tekin, B., Liang, K.J., Ma, H., Wang, W., Chen, X., Gleize, P., Xue, H., Lyu, S., et al.: Hoigpt: Learning long-sequence hand-object interaction with language models. In: Proceedings of the Computer Vision and Pattern Recognition Conference. pp. 7136–7146 (2025) [4](#)
 18. Jiang, S., Ye, Q., Xie, R., Huo, Y., Li, X., Zhou, Y., Chen, J.: In-hand 3d object reconstruction from a monocular rgb video. In: Proceedings of the AAAI Conference on Artificial Intelligence. vol. 38, pp. 2525–2533 (2024) [2](#)
 19. Kuang, Z., Ding, C., Yao, H.: Learning context with priors for 3d interacting hand-object pose estimation. In: Proceedings of the ACM International Conference on Multimedia. pp. 768–777 (2024) [2](#), [12](#)
 20. Lin, X., Zhu, M., Dang, R., Zhou, G., Shu, S., Lin, F., Liu, C., Chen, Q.: Clipose: Category-level object pose estimation with pre-trained vision-language knowledge. IEEE Transactions on Circuits and Systems for Video Technology **34**(10), 9125–9138 (2024) [4](#)
 21. Lin, Z., Ding, C., Yao, H., Kuang, Z., Huang, S.: Harmonious feature learning for interactive hand-object pose estimation. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 12989–12998 (2023) [2](#), [4](#), [10](#), [11](#), [12](#)
 22. Liu, J., Sun, W., Liu, C., Zhang, X., Fu, Q.: Robotic continuous grasping system by shape transformer-guided multi-object category-level 6d pose estimation. IEEE Transactions on Industrial Informatics **19**(11), 11171–11181 (2023) [1](#)
 23. Liu, J., Sun, W., Yang, H., Deng, P., Liu, C., Sebe, N., Rahmani, H., Mian, A.: Diff9d: Diffusion-based domain-generalized category-level 9-dof object pose estimation. IEEE Transactions on Pattern Analysis and Machine Intelligence **47**(7), 5520–5537 (2025) [1](#)

24. Liu, J., Sun, W., Yang, H., Zeng, Z., Liu, C., Zheng, J., Liu, X., Rahmani, H., Sebe, N., Mian, A.: Deep learning-based object pose estimation: A comprehensive survey. *International Journal of Computer Vision* **134**(2), 81 (2026) [4](#)
25. Liu, J., Sun, W., Zeng, K., Zheng, J., Yang, H., Rahmani, H., Mian, A., Wang, L.: Scalable unseen objects 6-dof absolute pose estimation with robotic integration. *IEEE Transactions on Robotics* **42**, 1884–1901 (2026) [1](#)
26. Liu, S., Jiang, H., Xu, J., Liu, S., Wang, X.: Semi-supervised 3d hand-object poses estimation with interactions in time. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 14687–14697 (2021) [3](#), [10](#), [12](#)
27. Liu, Y., Long, X., Yang, Z., Liu, Y., Habermann, M., Theobalt, C., Ma, Y., Wang, W.: Easyhoi: Unleashing the power of large models for reconstructing hand-object interactions in the wild. In: *Proceedings of the Computer Vision and Pattern Recognition Conference*. pp. 7037–7047 (2025) [4](#)
28. Park, J., Oh, Y., Moon, G., Choi, H., Lee, K.M.: Handoccnnet: Occlusion-robust 3d hand mesh estimation network. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 1496–1505 (2022) [10](#)
29. Qi, H., Zhao, C., Salzmänn, M., Mathis, A.: HoisdF: Constraining 3d hand-object pose estimation with global signed distance fields. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 10392–10402 (2024) [3](#), [6](#), [9](#), [10](#), [11](#), [12](#)
30. Romero, J., Tzionas, D., Black, M.J.: Embodied hands: Modeling and capturing hands and bodies together. *ACM Transactions on Graphics* **36**(6) (2017) [3](#)
31. Wang, R., Mao, W., Li, H.: Interacting hand-object pose estimation via dense mutual attention. In: *Proceedings of the IEEE/CVF winter conference on applications of computer vision*. pp. 5735–5745 (2023) [2](#)
32. Wang, Y., Xu, H., Heng, P.A., Fu, C.W.: Unihope: A unified approach for hand-only and hand-object pose estimation. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 12231–12241 (2025) [2](#), [10](#), [11](#), [12](#)
33. Wang, Z., Chen, J., Chen, Z., Xie, P., Chen, R., Yi, L.: Genh2r: Learning generalizable human-to-robot handover via scalable simulation demonstration and imitation. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 16362–16372 (2024) [1](#)
34. Xiang, Y., Schmidt, T., Narayanan, V., Fox, D.: Posecnn: A convolutional neural network for 6d object pose estimation in cluttered scenes. *arXiv preprint arXiv:1711.00199* (2017) [9](#)
35. Xu, H., Li, H., Wang, Y., Liu, S., Fu, C.W.: Handbooster: Boosting 3d hand-mesh reconstruction by conditional synthesis and sampling of hand-object interactions. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 10159–10169 (2024) [1](#)
36. Xu, H., Wang, T., Tang, X., Fu, C.W.: H2onet: Hand-occlusion-and-orientation-aware network for real-time 3d hand mesh reconstruction. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 17048–17058 (2023) [10](#)
37. Yang, H., Sun, W., Liu, J., Zheng, J., Dai, Z., Mian, A.: LancopE: Language-guided category-level object pose estimation from a single rgb image. *IEEE Robotics and Automation Letters* (2025) [4](#)
38. Yang, H., Sun, W., Liu, J., Zheng, J., Xiao, J., Mian, A.: Occlusion-aware 3d hand-object pose estimation with masked autoencoders. *arXiv preprint arXiv:2506.10816* (2025) [6](#)

39. Yang, H., Sun, W., Liu, J., Zheng, J., Zeng, Z., Mian, A.: Rgb-based category-level object pose estimation via depth recovery and adaptive refinement. *IEEE Robotics and Automation Letters* (2025) [4](#)
40. Ye, Y., Gupta, A., Kitani, K., Tulsiani, S.: G-hop: Generative hand-object prior for interaction reconstruction and grasp synthesis. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 1911–1920 (2024) [4](#)
41. Ye, Y., Li, X., Gupta, A., De Mello, S., Birchfield, S., Song, J., Tulsiani, S., Liu, S.: Affordance diffusion: Synthesizing hand-object interactions. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 22479–22489 (2023) [1](#)
42. Zhang, C., Di, Y., Zhang, R., Zhai, G., Manhardt, F., Tombari, F., Ji, X.: Ddf-ho: Hand-held object reconstruction via conditional directed distance field. *Advances in Neural Information Processing Systems* **36**, 56871–56884 (2023) [3](#)
43. Zhang, C., Jiao, G., Di, Y., Wang, G., Huang, Z., Zhang, R., Manhardt, F., Fu, B., Tombari, F., Ji, X.: Moho: Learning single-view hand-held object reconstruction with multi-view occlusion-aware supervision. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 9992–10002 (2024) [2](#)
44. Zhang, H., Christen, S., Fan, Z., Hilliges, O., Song, J.: Grasppl: Generating grasping motions for diverse objects at scale. In: *European Conference on Computer Vision*. pp. 386–403 (2024) [1](#)
45. Zhang, M., Fu, Y., Ding, Z., Liu, S., Tu, Z., Wang, X.: Hoidiffusion: Generating realistic 3d hand-object interaction data. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 8521–8531 (2024) [4](#)
46. Zhou, Z., Zhou, S., Lv, Z., Zou, M., Tang, Y., Liang, J.: A simple baseline for efficient hand mesh reconstruction. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 1367–1376 (2024) [10](#)
47. Zimmermann, C., Brox, T.: Learning to estimate 3d hand pose from single rgb images. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision*. pp. 4903–4911 (2017) [10](#)
48. Zimmermann, C., Ceylan, D., Yang, J., Russell, B., Argus, M., Brox, T.: Freihand: A dataset for markerless capture of hand pose and shape from single rgb images. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision*. pp. 813–822 (2019) [10](#)