
Learning Dense Hand Contact Estimation from Imbalanced Data

Daniel Sungho Jung¹ Kyoung Mu Lee^{1,2}

¹IPAI, ²Dept. of ECE & ASRI, Seoul National University, Korea
{dqj5182, kyoungmu}@snu.ac.kr

Abstract

Hands are essential to human interaction, and exploring contact between hands and the world can promote comprehensive understanding of their function. Recently, there have been growing number of hand interaction datasets that cover interaction with object, other hand, scene, and body. Despite the significance of the task and increasing high-quality data, how to effectively learn dense hand contact estimation remains largely underexplored. There are two major challenges for learning dense hand contact estimation. First, there exists class imbalance issue from hand contact datasets where majority of regions are not in contact. Second, hand contact datasets contain spatial imbalance issue with most of hand contact exhibited in finger tips, resulting in challenges for generalization towards contacts in other hand regions. To tackle these issues, we present a framework that learns dense HAnd COntact estimation (HACO) from imbalanced data. To resolve the class imbalance issue, we introduce balanced contact sampling, which builds and samples from multiple sampling groups that fairly represent diverse contact statistics for both contact and non-contact vertices. Moreover, to address the spatial imbalance issue, we propose vertex-level class-balanced (VCB) loss, which incorporates spatially varying contact distribution by separately reweighting loss contribution of each vertex based on its contact frequency across dataset. As a result, we effectively learn to predict dense hand contact estimation with large-scale hand contact data without suffering from class and spatial imbalance issue. The codes are available at https://github.com/dqj5182/HACO_RELEASE.

1 Introduction

From infancy, humans rely on physical contact to perceive and interact with the surrounding environment. Among the various mediums of contact, hands play a predominant role, facilitating essential tasks and enabling effective communication. Consequently, developing a robust hand contact estimation model is crucial for advancing our understanding of hand interactions and addressing various challenges [14, 35] that require accurate hand contact estimation.

Over the recent years, we have witnessed significant advances in large-scale interaction datasets involving hands. Hand-object interaction datasets [19, 5, 4, 15, 16, 12, 36, 24] focused on capturing hand grasps of objects. Hand-hand interaction datasets [59, 41] considered two hand interaction from a single person. Hand-face interaction dataset [55] investigated on interaction between hand and deformable face. Hand-scene interaction datasets [18, 22, 58] extended the scope of interaction to include environments featuring the ground, walls, and large-scale objects. Hand-body interaction dataset [65] covered broad range of interaction with human body as part of human-human interaction. Despite the wide range of diversity in the hand interaction datasets with contact annotations [19, 5, 4, 15, 16, 12, 36, 24, 41, 59, 55, 17, 22, 65, 62], to the best of our knowledge, there has been no attempt to build hand contact estimation model trained on such diverse hand contact. Based on such diverse sets of hand contact data, we aim to build an effective dense hand contact estimation model.

There are two major issues for training a dense hand contact estimation. First, there exists class imbalance between contact and non-contact class. Human participants of the hand interaction datasets are often instructed to perform specific actions such as picking up an object [5, 15, 36] and pointing to specific regions of objects [24] or other hands [41]. As most of such actions require precise control with interaction by a small region on hand surface, most of the hand regions are overlooked and excluded from being in contact. In Figure 1a, we provide comparison on the occurrence of contact and non-contact in major hand contact datasets for all vertices of MANO hand model [54]. This observation shows severe class imbalance issue of hand contact from existing hand contact datasets. Numerically, DexYCB [5] dataset has roughly 2.7:1 ratio of non-contact to contact vertices. More severe InterHand2.6M [41] dataset has huge imbalance ratio of 19.5:1. Such imbalance is also shown in Decaf [55] dataset with 21.7:1 ratio. As stated by numerous studies on data imbalance [3, 6, 52, 26], such data imbalance causes poor performance on underrepresented classes, which is in fact the contacting cases in our scenario. Second, the spatial imbalance issue of hand contact data exists throughout the majority of hand contact datasets. This imbalance arises from the fact that hand interaction predominantly involves the fingertips, which provide a high degree of freedom in both movement and rotation. Due to such high maneuverability and the precise nature of actions from motion capture datasets, most of hand contact from the datasets are significantly skewed to finger tip regions. As shown in Figure 1b, the heatmap for hand contact from hand contact datasets are highly concentrated at the finger tips. Due to such spatial imbalance issue of hand contact datasets, models directly trained on the datasets struggle to generalize on diverse contact patterns across hand.

Therefore, we propose HACO, a framework that learns dense hand contact estimation from imbalanced data. Initially, HACO is trained on a large-scale, assorted dataset composed of 14 datasets, shown in Table 1, that contains various hand interactions to leverage the power of large-scale training [2, 53, 23, 64, 49]. However, this large-scale dataset still exhibits inherent imbalance in terms of both class and spatial distribution. To address these issues, we introduce two strategies designed to fully exploit the potential of such large-scale data. First, we introduce *Balanced Contact Sampling (BCS)*, which mitigates class imbalance in hand contact datasets by dividing the dataset into multiple sampling groups, each representing different statistics of contact and non-contact. To ensure fair representation of the contact statistics, we compute a contact balance score that measures how much each hand contact instance deviates from the dataset-wide average. This metric enables us to assess whether each instance is under-represented or over-represented within the original dataset and to adjust the sampling ratios accordingly, ensuring all contact types are well represented. Second, we present the *Vertex-Level Class-Balanced (VCB) loss*, which applies spatially varying weights to each vertex to address the spatial imbalance issue in hand contact datasets. Inspired by the core idea of class-balanced (CB) loss [6], which re-weights the loss based on class frequency, we adapt this concept to the spatial domain by computing a separate loss weight for each vertex based on how well contact and non-contact are represented at the vertex within the dataset. While the original CB loss applies global class-level weighting, our method performs fine-grained, vertex-level weighting, allowing the model to focus on under-represented contact patterns in a spatially-aware manner.

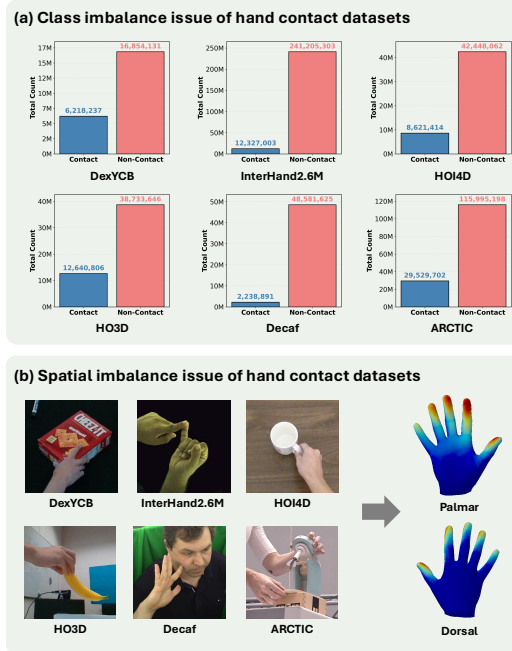


Figure 1: **Two challenges for dense hand contact estimation in the wild.** First, hand contact datasets suffer from class imbalance, as the majority of vertices contain no contact. Second, spatial imbalance arises because contact points are predominantly concentrated at the fingertips. Due to these issues, models trained on such data struggle to generalize to diverse contact patterns across the hand.

Table 1: **Dataset configuration.** We leverage 14 datasets with various hand interaction.

Dataset	Interaction	Domain	# of images	# of subjects	# of objects
ObMan [19]	Rigid Object	Synthetic	152K	20	8
DexYCB [5]	Rigid Object	Real / Indoor	582K	10	20
MOW [4]	Rigid Object	Real / Outdoor	0.5K	450	121
HO3D [15]	Rigid Object	Real / Indoor	78K	10	10
H2O3D [16]	Rigid Object + Hand	Real / Indoor	76K	5	10
ARCTIC [12]	Articulated Object + Hand	Real / Indoor / Studio	2.1M	10	11
HOI4D [36]	Rigid Object + Hand	Real / Indoor	2.4M	4	16
H2O [24]	Rigid Object + Hand	Real / Indoor	571K	4	8
InterHand2.6M [41]	Hand	Real / Indoor / Studio	2.6M	27	–
HIC [59]	Hand	Real / Indoor	36K	1	–
PROX [17]	Scene	Real / Indoor / Studio	100K	20	–
RICH [22]	Scene	Real / Outdoor	540K	22	–
Decaf [55]	Face (Whole-Body)	Real / Indoor	100K	8	–
Hi4D [65]	Whole-Body	Real / Indoor / Studio	11K	40	–

As a result, HACO achieves state-of-the-art performance across diverse hand contact scenarios, including hand-object, hand-hand, hand-scene, and hand-body interactions, consistently outperforming existing methods.

Our key contributions are as follows:

- We introduce HACO, a novel framework that addresses the data imbalance problems in hand contact datasets and enables effective learning of dense hand contact estimation from large-scale data.
- To mitigate class imbalance in existing hand contact datasets, we propose Balanced Contact Sampling (BCS), which constructs multiple sampling groups to fairly represent diverse contact statistics.
- To handle spatial imbalance, we present the Vertex-Level Class-Balanced (VCB) loss, which computes a vertex-specific weighting factor based on the dataset-wide contact distribution.
- In the end, HACO, the first method trained for dense hand contact estimation on a large-scale dataset, achieves substantial performance gains across diverse contact scenarios.

2 Related works

Dense hand contact estimation. Most of the existing methods that explore dense hand contact estimation are devoted to dense human-scene contact estimation [18, 22, 58, 43]. Given an image of human-scene interaction, such methods estimate contactness for each vertex of 3D human model (e.g., SMPL [37], SMPL-X [48]). POSA [18] is based upon the idea that the human-scene contact is largely influenced by the 3D pose of human. Hence, the work employs a conditional variational autoencoder framework [56] to learn human-scene contact with the 3D vertex position of posed human as condition. BSTRO [22] takes another approach for estimating human-scene contact, which directly relies on visual input with Transformer-based architecture [9]. To train the model, they proposed RICH dataset [22] that annotates dense human-scene contact by capturing 3D geometry of body and scene with multi-view laser scanner in both indoor and outdoor environment. Lastly, to improve generalization ability of human-scene contact estimation, DECO [58] focused on how to annotate human-scene contact in the wild. They introduced simple yet effective contact annotation method that simply asks crowd-sourced annotators to paint 3D human model (i.e., SMPL [37]) in neutral T-pose, which significantly enhanced the potential of scaling contact data. Despite numerous attempts on dense human-scene contact estimation, there lacks in-depth investigation on dense hand contact estimation. Furthermore, we found that dense hand contact estimation has unique challenges compared to human-scene contact estimation. First, due to class imbalance of hand contact datasets, the distribution of the datasets are often skewed towards non-contact class, making it non-trivial to train hand contact estimation from the datasets. Second, there exists spatial imbalance issue of hand contact datasets where majority of hand contact occurs in finger tips with insufficient data for hand contact in rest of the regions (e.g., dorsum). We introduce HACO, a dense hand contact estimation method that effectively tackles such data imbalance issues of hand contact datasets.

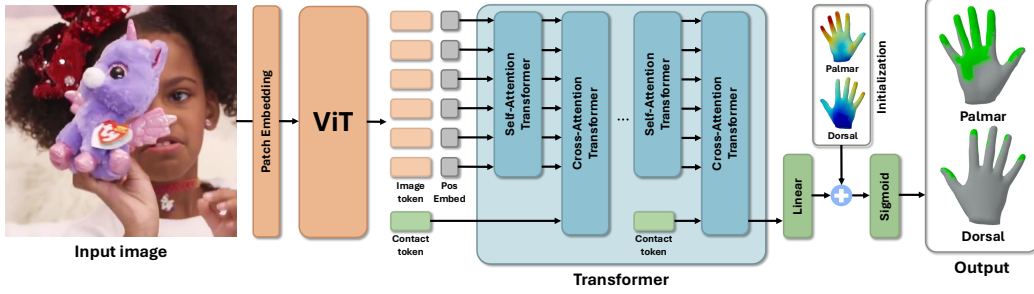


Figure 2: **Overall pipeline of HACO.** Our method encodes input image as image tokens with a ViT backbone after patch embedding layers. Given the image tokens along with positional embeddings and a contact token, multiple layers of self-attention Transformer and cross-attention Transformer produce an output token. Lastly, the output token is passed through a linear layer and combined with the contact initialization, followed by a sigmoid activation to output the final hand contact prediction.

3D hand and object interaction reasoning. Dense hand contact reveals key spatial and semantic relationship between 3D hand and object. Initially, dense hand contact was studied in the field of grasping with specific tasks such as grasp contact prediction [1], contact-guided hand and object pose refinement [14], and contact-conditioned grasp generation [34]. Such works made tremendous impact on robotics, especially on how robots can naturally grasp an object. However, their dense hand contact were mostly predicted from 3D geometry of hand and object without exploring how to directly predict the dense hand contact from visual input. Recently, the field of 3D hand and object reconstruction showed potential on the effectiveness of contact for improving both reconstruction accuracy and natural interaction between hand and object. ContactOpt [14] proposed a differentiable contact optimization that optimizes hand and object poses inferred from image-based methods with estimated contact to achieve natural interaction. EasyHOI [35] introduced a prior-guided hand-object interaction optimizer that refines initial 3D hand and object using heuristically determined contact regions to improve interaction quality. Our HACO is applicable to these downstream tasks as long as visual input is available. To demonstrate HACO’s effectiveness, we show that our predicted dense contact improves 3D grasp optimization over geometry-driven contact estimation from ContactOpt [14], and 3D reconstruction over the previously used heuristic contact regions in EasyHOI [35].

3 Method

3.1 Model architecture

Given an RGB image $\mathbf{I} \in \mathbb{R}^{3 \times H \times W}$, where H and W denote the height and width of the image, we first embed the image into multiple tokens with a patch embedding layer and process these tokens using a Vision Transformer (ViT) [10]. Following ViTPose [63] and HaMeR [49], we reshape the tokens and formulate the image feature $\mathbf{F} \in \mathbb{R}^{1280 \times 16 \times 12}$. Additionally, we build a contact token, which acts as a query token for hand contact estimation. To inject image feature $\mathbf{F}$ into the contact token, we process the image feature $\mathbf{F}$ and positional embedding with consecutive self-attention Transformer and cross-attention Transformer with contact token as query and image feature $\mathbf{F}$ as key-value pair. Afterwards, the output contact token is processed with linear layer. In order to resemble residual layers [20] that stabilize training, we add contact initialization to the output of linear layer. The contact initialization is a learnable embedding that naturally learns the most effective initial contact estimation during training. Lastly, a sigmoid layer converts the final contact logits $\mathbf{C} \in \mathbb{R}^V$ into the estimated dense hand contact probabilities, where $V = 778$ is MANO hand vertices [54].

3.2 Balanced contact sampling

A hand contact dataset $\mathbf{D} = \{\mathbf{H}_i = (\mathbf{v}_i, \mathbf{c}_i)\}_{i=1}^N$ consists of N hand instances, where each $\mathbf{H}_i$ contains the 3D hand vertices $\mathbf{v}_i \in \mathbb{R}^{V \times 3}$ and a binary vertex contact vector $\mathbf{c}_i \in \{0, 1\}^V$. As illustrated in Figure 1a, the dataset $\mathbf{D}$ exhibits a strong class imbalance, with many hand instances lacking contact. To alleviate this issue, we propose balanced contact sampling (BCS), which constructs sampling bins $\mathbf{B} = \{\mathcal{B}_1, \mathcal{B}_2, \dots, \mathcal{B}_K\}$ consisting of K groups of hand instances. Each bin $\mathcal{B}_k$ is formed by binning samples based on their contact balance scores s_i , which reflect how each hand’s

contact $\mathbf{c}_i$ deviates from the dataset-wide average $\bar{\mathbf{c}}$. To compute s_i , we first define the dataset-wide contact mean as $\bar{\mathbf{c}} = \frac{1}{N} \sum_{i=1}^N \mathbf{c}_i$, which represents the contact probability at each vertex across all samples. The contact balance score s_i for each sample is then defined as:

$$s_i = \frac{1}{V} (\mathbf{c}_i^\top (1 - \bar{\mathbf{c}}) - \mathbf{c}_i^\top \bar{\mathbf{c}}) \quad (1)$$

where higher scores correspond to hand instances $\mathbf{H}_i$ whose contact patterns deviate more from the dataset-wide average, enabling balanced grouping across diverse contact statistics that span both contact and non-contact samples. Based on the distribution of contact balance scores $\{s_i\}_{i=1}^N$, we partition the dataset into K bins $\mathcal{B}_1, \dots, \mathcal{B}_K$ by applying non-linear binning. We first compute the minimum and maximum contact balance scores as $s_{\min} = \min_i s_i$ and $s_{\max} = \max_i s_i$, and define $K+1$ bin edges using a logarithmic spacing scheme controlled by a curvature parameter $\beta > 0$ (e.g., $\beta = 5$):

$$\tau_k = s_{\min} + (s_{\max} - s_{\min}) \cdot \frac{\log(1 + \beta \cdot x_k)}{\log(1 + \beta)}, \quad x_k = \frac{k}{K}, \quad k = 0, \dots, K. \quad (2)$$

This results in bin edges $\{\tau_0, \tau_1, \dots, \tau_K\}$ that allocate finer resolution to higher contact scores. Each group $\mathcal{B}_k$ is then constructed by selecting hand instances whose contact balance scores fall within the corresponding interval:

$$\mathcal{B}_k = \{\mathbf{H}_i \in \mathbf{D} \mid s_i \in [\tau_{k-1}, \tau_k)\} \quad (3)$$

Due to the skewed distribution of scores s_i , the bins $\mathcal{B}_k$ may differ in size. To equalize their contribution, we apply stratified resampling to ensure each bin contains the same number of hand instances in the final training set. In the end, this binning strategy enables more accurate grouping and fair sampling of contact-rich hand instances that are underrepresented in the original dataset $\mathbf{D}$, improving the representational balance of the sampling bins $\mathbf{B}$.

3.3 Vertex-level class-balanced loss

As illustrated in Figure 1b, hand contact datasets exhibit severe spatial imbalance, with contact heavily concentrated at the fingertips. To address this issue, we propose a vertex-level class-balanced (VCB) loss, which builds upon the class-balanced (CB) sigmoid binary cross-entropy loss [6]. We first formulate hand contact estimation as a binary classification problem. A simple binary cross-entropy (BCE) loss is defined as $\ell_{\text{BCE}}(y, p) = -y \log(p) - (1 - y) \log(1 - p)$, where $y \in \{0, 1\}$ is the ground-truth label and $p \in [0, 1]$ is the predicted probability. When applied over all hand vertices V , the overall BCE loss becomes:

$$\mathcal{L}_{\text{BCE}} = \frac{1}{|V|} \sum_{v \in V} \ell_{\text{BCE}}(y_v, p_v), \quad (4)$$

where we denote ground-truth contact for vertex v as $y_v \in \{0, 1\}$ and predicted contact probability for vertex v as $p_v = \sigma(z)$ with sigmoid function as σ and predicted logit as z .

From the BCE loss, class-balanced (CB) loss [6] introduces a weighting factor α_c for each class c , assigning greater importance to under-represented classes. In our case, the class c corresponds to the binary contact label $y \in \{0, 1\}$. Accordingly, we denote the class-balanced weighting factor as α_y . Simply, we can write CB loss as follows:

$$\mathcal{L}_{\text{CB}} = \frac{1}{|V|} \sum_{v \in V} \alpha_y \ell_{\text{BCE}}(y_v, p_v). \quad (5)$$

Specifically, the weighting factor α_y is defined using the effective number of samples $E_n^{(y)}$, where

$$\alpha_y = \frac{1}{E_n^{(y)}} = \frac{1 - \beta}{1 - \beta^{n_y}}. \quad (6)$$

Here, n_y denotes the number of occurrences of class y , and $\beta \in [0, 1)$ is a hyperparameter that controls the influence of sample count on the weighting. When the contact class y appears frequently in the dataset, the contact frequency value n_y becomes large, which increases the effective number $E_n^{(y)} = \frac{1 - \beta^{n_y}}{1 - \beta}$ and results in a smaller class-balanced weight $\alpha_y = \frac{1}{E_n^{(y)}}$. As a result, the CB loss down-weights the contribution of frequent classes and helps mitigate class imbalance in dense hand contact datasets.

Nevertheless, CB loss [6] is defined at the class level, based on the binary contact label $\mathbf{y} \in \{0, 1\}$. If we directly apply CB loss, the weighting factor α_y takes only two distinct values, one for contact and one for non-contact, and is applied uniformly across all spatial regions. In the DexYCB dataset [5], contact and non-contact vertices account for approximately 76.86% and 23.14% of all vertex labels, respectively. Accordingly, the weighting factor α_y is computed using sample counts $n_y = (0.7686 \times n, 0.2314 \times n)$, where n is the total number of contact labels in the dataset. However, this count n_y is applied uniformly across all vertices, without considering spatial variation. For example, fingertip vertices are much more likely to be in contact, while dorsal hand vertices are rarely in contact. As a result, the class-balanced weight α_y and effective sample number $E_n^{(y)}$ do not accurately reflect the local distribution of contact and non-contact at each vertex. To address this limitation, we extend the CB loss formulation to operate at the vertex level.

We introduce the vertex-level class-balanced (VCB) loss, which computes a spatially varying weighting factor α_{y_v} and effective number of samples $E_n^{(y_v)}$ for each vertex v . To account for the spatial imbalance in hand contact datasets, we simply assign a separate weighting factor α_{y_v} to each vertex, depending on the local distribution of contact labels. Given the number of samples n_{y_v} for contact class y at vertex v , the weighting factor is computed as:

$$\alpha_{y_v} = \frac{1}{E_n^{(y_v)}} = \frac{1 - \beta}{1 - \beta^{n_{y_v}}}. \quad (7)$$

Using this spatially varying weight, the vertex-level class-balanced loss is defined as:

$$\mathcal{L}_{\text{VCB}} = \frac{1}{|V|} \sum_{v \in V} \alpha_{y_v} \ell_{\text{BCE}}(y_v, p_v). \quad (8)$$

This spatial adaptivity enables the VCB loss to effectively mitigate the spatial imbalance inherent in hand contact datasets. The VCB loss is applied through a progressive weighting strategy that combines the CB loss and VCB loss throughout training. At the beginning of training, only the CB loss is applied, while the weight of the VCB component increases linearly with each epoch. The VCB weight reaches a fixed maximum at the final epoch, resulting in a smooth transition from global to vertex-adaptive supervision. Lastly, we include a lightweight false-positive regularization term to suppress overly confident contact predictions for non-contact vertices.

3.4 Final outputs and loss functions

The final dense hand contact output supports multi-level supervision through projections to coarser resolutions. From the dense contact logits $\mathbf{C}$ regressed by the final linear layer, we build multi-level hand contact predictions $\mathbf{C}' = \{\mathbf{C}_i = \mathcal{J}_i^{v_i} \times \mathbf{C}\}_{i=1}^N$ for $N = 4$ and $\mathbf{v}_i \in \{778, 336, 84, 21\}$ where $\mathcal{J}_i^{v_i} \in \mathbb{R}^{v_i \times V}$ is a regressor that maps full MANO vertices into coarse representations. We then apply a sigmoid layer to convert the logits into contact probabilities. To train HACO, we apply three loss functions: the vertex-level class-balanced (VCB) loss for multi-level hand contact supervision, a regularization loss, and a smoothness loss. Please refer to Section 3.3 for details on the VCB loss, and to the Appendix for the regularization and smoothness losses. Briefly, the regularization loss is defined as the L1 loss between the dense contact prediction and the dataset-wide mean ground-truth contact, encouraging the prediction to stay close to the average. The smoothness loss measures how predicted contact and non-contact regions are spatially disconnected. This encourages HACO to predict a small number of large contact regions rather than many small, fragmented ones. We set the weights of the VCB loss, regularization loss, and smoothness loss to 1, 0.1, and 1, respectively.

4 Implementation details

PyTorch [47] is used for implementation. Our backbone is initialized with the pre-trained weights of publicly released HaMeR [49]. Following HaMeR, we apply data augmentations including random scaling, cropping, and rotation. To improve robustness to degraded inputs, we additionally perform low-resolution, noise, and blur augmentations. We use the AdamW optimizer [38] with a learning rate of 10^{-5} and a mini-batch size of 24. For stable convergence, the learning rate is reduced by a factor of 0.9 after 5 and 10 epochs. We train HACO for 10 epochs on a single NVIDIA A6000 GPU.

5 Experiments

5.1 Datasets

We select 14 datasets with diverse hand interactions including ObMan [19], DexYCB [5], MOW [4], HO3D [15], H2O3D [16], ARCTIC [12], HOI4D [36], H2O [24] for hand-object interaction, InterHand2.6M [41], HIC [59] for hand-hand interaction, PROX [17], RICH [22] for hand-scene interaction, and Decaf [55], Hi4D [65] for hand-body interaction. To reduce redundancy from large video datasets, we employ sampling ratio of 5, 10, 5 for HOI4D [36], InterHand2.6M [41], and Decaf [55] dataset, respectively. Note that this is separate from the balanced contact sampling in Section 3.2. For ARCTIC [12] dataset, we select data captured from an egocentric viewpoint, while for HOI4D [36], we use data involving rigid objects due to their uniqueness. Since the 3D hand annotations in PROX [17] are inaccurate in the qualitative split, we use only the quantitative subset. For the RICH dataset, we follow the official split used in BSTRO [22] for fair comparison. In total, this results in 655K images with ground-truth dense hand contact labels. For details on our method on extracting ground-truth dense hand contact labels, please refer to Appendix.

5.2 Evaluation metrics

To evaluate dense hand contact estimation, we compute precision, recall, and F1-score. However, recall and F1-score are undefined for fully non-contact samples. Thus, we skip fully non-contacting hands during evaluation. We also assess contact quality on two downstream tasks: 3D hand grasp optimization [14] and 3D hand-object reconstruction [35]. For 3D hand grasp optimization [14], we evaluate Intersection Volume (Inter Vol.), Mean Per-Joint Position Error (MPJPE), and contact metrics. For 3D hand-object reconstruction, we compare HACO using Mean Per-Vertex Position Error (MPVPE), MPJPE, CD_{ho} , $F-5_{ho}$, and $F-10_{ho}$ against the contact method of EasyHOI. $F-5_{ho}$ and $F-10_{ho}$ denote F-scores of reconstructed 3D hand and object using 5mm and 10mm thresholds.

5.3 Ablation study

Effectiveness of balanced contact sampling. In Table 2, the proposed balanced contact sampling strategy significantly enhances the performance across all metrics. Compared to the model trained without sampling, our method improves 1.0% in precision, 12.0% in recall, and 8.5% in F1-score.

The substantial boost in recall, in particular, highlights the effectiveness of balanced sampling in mitigating class imbalance by exposing the model to more positive contact instances. Overall, these results demonstrate that our sampling strategy plays a critical role in improving our model’s ability to detect contact accurately and comprehensively.

Table 2: **Ablation on balanced contact sampling strategy on MOW [4] dataset.**

Methods	Precision↑	Recall↑	F1-Score↑
w/o sampling	0.520	0.542	0.481
w/ sampling (Ours)	0.525	0.607	0.522

Effectiveness of vertex-level class-balanced loss. In Table 3, the proposed VCB loss achieves the best overall performance by effectively addressing data imbalance at a finer granularity.

Unlike CB loss, which applies class balancing at the hand-level by weighting positive and negative samples across all vertices, VCB loss operates at the per-vertex level, allowing the model to learn contact patterns from individual vertex more accurately. This vertex-wise weighting enables the model to better capture subtle and spatially varying contact signals, especially in under-represented regions, leading to improved recall without sacrificing precision. As a result, VCB loss yields the highest F1-score among all loss-based methods, demonstrating that resolving class imbalance locally, rather than globally, is key to effective dense contact estimation.

Table 3: **Comparison of various loss-based techniques for data imbalance on MOW [4] dataset.**

Methods	Precision↑	Recall↑	F1-Score↑
CE loss	0.530	0.294	0.348
L1 loss	0.521	0.392	0.413
L2 loss	<u>0.531</u>	0.298	0.352
Focal loss [32]	0.518	0.387	0.409
CB loss [6]	0.484	<u>0.534</u>	<u>0.465</u>
CB Focal loss [6]	0.522	0.360	0.392
LDAM loss [3]	0.532	0.224	0.293
Asymmetric loss [52]	0.484	0.479	0.440
Poly loss [26]	0.528	0.324	0.371
VCB loss (Ours)	0.525	0.607	0.522

Table 4: **Comparison of various training dataset configurations on MOW [4], HIC [59], RICH [22], Hi4D [65].** HO: Hand-object, HH: Hand-hand, HS: Hand-scene, HB: Hand-body.

Test dataset	Train dataset	Precision $\uparrow$	Recall $\uparrow$	F1-Score $\uparrow$
MOW [4]	HS	0.397	0.516	<u>0.411</u>
	HS+HH	0.405	0.420	0.373
	HS+HH+HB	<u>0.425</u>	<u>0.551</u>	0.408
	HS+HH+HB+HO (Ours)	0.525	0.607	0.522
HIC [59]	HS	<u>0.188</u>	0.231	<u>0.187</u>
	HS+HO	0.094	0.260	0.134
	HS+HO+HB	0.114	<u>0.265</u>	0.158
	HS+HO+HB+HH (Ours)	0.216	0.409	0.263
RICH [22]	HO	0.551	0.547	<u>0.492</u>
	HO+HH	<u>0.554</u>	0.522	0.471
	HO+HH+HB	<u>0.554</u>	<u>0.581</u>	0.490
	HO+HH+HB+HS (Ours)	0.741	0.899	0.781
Hi4D [65]	HS	0.515	0.439	0.419
	HS+HO	0.610	0.503	0.496
	HS+HO+HH	<u>0.596</u>	<u>0.513</u>	<u>0.498</u>
	HS+HO+HH+HB (Ours)	0.555	0.636	0.565

Effectiveness of large-scale data configuration. Our final model benefits significantly from the inclusion of all four interaction types: hand-object (HO), hand-hand (HH), hand-body (HB), and hand-scene (HS). As shown in Table 4, performance improves consistently across all evaluated datasets as more interaction types are incrementally incorporated, with the full combination achieving the highest overall F1-score. Removing HO leads to a substantial drop in recall, further highlighting its essential role in capturing diverse contact patterns. Although HB and HH alone contribute modestly, their addition enhances F1-scores and stabilizes overall performance, indicating their complementary effects. Omitting HS results in reduced F1-score despite strong contributions from the other components, suggesting its importance in separating hand contacts from background context. Ultimately, only the full configuration achieves strong and balanced performance across all metrics and datasets, demonstrating that each interaction type contributes uniquely and that leveraging the full diversity is crucial for learning accurate and generalizable contact representations.

5.4 Comparison with state-of-the-art methods

Dense hand contact estimation. Table 5 presents a comparison between our HACO and state-of-the-art approaches of POSA [18], BSTRO [22], and DECO [58] on the MOW [4] dataset.

Our method achieves substantial improvements across all metrics, significantly outperforming the strongest prior method, DECO [58]. POSA [18] shows limited performance due to its sole reliance on pose priors [13] without image evidence, while BSTRO [22] suffers from the lack of large-scale, diverse contact data and does not address class or spatial imbalance during training. DECO [58] improved over earlier methods, but performance remains constrained by the lack of diverse supervision and mechanisms for handling imbalance. In contrast, HACO addresses these challenges by training on diverse contact configurations, initializing contact prediction with a learned prior, and tackling class and spatial imbalance through targeted loss design. Figure 3 further highlights HACO’s qualitative superiority, showing more precise and anatomically plausible contact patterns across scenarios, such as fingertip and palmar contact when holding a mic, partial palmar contact when dicing with knife, and thumb–index finger contact during grasp of a pencil. Overall, our method shows superior performance over prior SOTA methods, demonstrating its potential as a strong model for dense hand contact estimation by addressing key challenges that limit earlier approaches.

Table 5: **Comparison with SOTA methods of hand contact estimation on MOW [4] dataset.**

Methods	Precision $\uparrow$	Recall $\uparrow$	F1-Score $\uparrow$
POSA [18]	0.134	0.128	0.101
BSTRO [22]	0.204	0.126	0.112
DECO [58]	<u>0.246</u>	<u>0.235</u>	<u>0.197</u>
HACO (Ours)	0.525	0.607	0.522

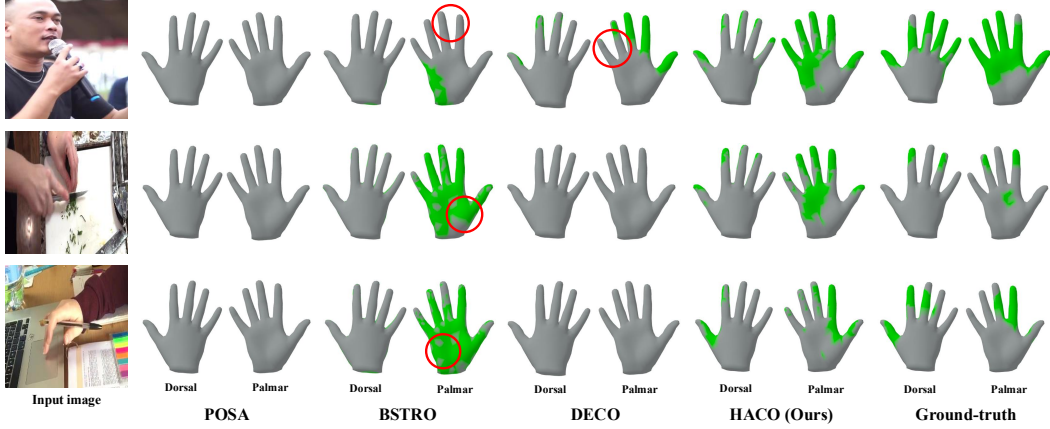


Figure 3: **Qualitative comparison of dense hand contact estimation with POSA [18], BSTRO [22], DECO [58] on MOW [4] dataset.** We highlight exemplar regions where HACO outperforms previous methods. Note that we only predict right hand contact.

3D hand grasp optimization. Table 6 compares our HACO with the prior state-of-the-art method, DeepContact [14], for contact-guided 3D grasp optimization using predicted hand-object poses from HFL-Net [33] on the DexYCB [5] dataset. Despite DeepContact having access to full 3D hand and object meshes for contact estimation, our method achieves comparable or superior results using only image input. HACO attains lower MPJPE, which indicates more accurate hand articulation. Furthermore, it achieves higher precision, recall, and F1-score, demonstrating improved contact quality and coverage. Overall, these results show that HACO enables more effective and robust contact-aware grasp optimization, even without access to full 3D hand and object geometry, highlighting its practical applicability in real-world scenarios.

Table 6: **Comparison of various contact estimation methods for 3D hand grasp optimization on DexYCB [5] dataset.**

Methods	Inter Vol.↓	MPJPE↓	Precision↑	Recall↑	F1-Score↑
DeepContact [14]	29.175	37.155	0.522	0.830	0.612
HACO (Ours)	29.264	36.520	0.558	0.877	0.666

3D hand and object reconstruction. Table 7 compares the performance of EasyHOI [35] between using its original contact estimation module and using HACO, keeping all other components unchanged. This isolates the impact of the contact on reconstruction performance. HACO outperforms the original across all metrics, including MPVPE, MPJPE, Chamfer Distance (CD_{ho}), and F-scores. These consistent gains suggest that HACO produces more accurate and physically plausible contact predictions, directly improving hand-object alignment and reconstruction quality. Figure 4 further supports this, showing that the 3D reconstruction with contact from HACO more precisely captures interaction between the hand and the pen, unlike EasyHOI’s heuristic-based approach.

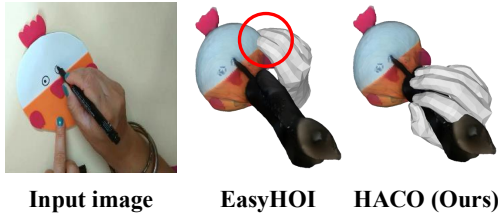


Figure 4: **Qualitative comparison of 3D hand-object reconstruction on MOW [4] dataset.**

Table 7: **Comparison of various contact estimation methods for 3D hand and object reconstruction with EasyHOI [35] on MOW [4] dataset.** PVE and PJE refer to MPVPE and MPJPE on vertices and joints between estimated and ground-truth 3D hand, respectively.

Methods	PVE↓	PJE↓	CD_{ho} ↓	$F-5_{ho}$ ↑	$F-10_{ho}$ ↑
EasyHOI [35]	21.254	20.973	8.338	0.120	0.230
HACO (Ours)	21.093	20.845	8.186	0.122	0.231

6 Conclusion

We propose HACO, a novel and powerful method that learns dense hand contact estimation from imbalanced data. For class imbalance and spatial imbalance in hand contact datasets, we propose balanced contact sampling and vertex-level class-balanced loss. As a result, our HACO outperforms previous methods by a significant margin on dense hand contact estimation and shows effectiveness on 3D grasp optimization and 3D hand and object reconstruction.

Acknowledgements

This work was supported in part by the IITP grants [No.2021-0-01343, No.2023-0-00156, and Artificial Intelligence Graduate School Program (Seoul National University) - No.2021-0-02068], and the Industrial Technology Alchemist Project [No.RS-2024-00432410] funded by MOTIE, Korea.

References

- [1] Samarth Brahmabhatt, Cusuh Ham, Charles C Kemp, and James Hays. ContactDB: Analyzing and predicting grasp contact via thermal imaging. In *CVPR*, 2019.
- [2] Tom B. Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu, Clemens Winter, Christopher Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. Language models are few-shot learners. In *NeurIPS*, 2020.
- [3] Kaidi Cao, Colin Wei, Adrien Gaidon, Nikos Arechiga, and Tengyu Ma. Learning imbalanced datasets with label-distribution-aware margin loss. In *NeurIPS*, 2019.
- [4] Zhe Cao, Ilija Radosavovic, Angjoo Kanazawa, and Jitendra Malik. Reconstructing hand-object interactions in the wild. In *ICCV*, 2021.
- [5] Yu-Wei Chao, Wei Yang, Yu Xiang, Pavlo Molchanov, Ankur Handa, Jonathan Tremblay, Yashraj S Narang, Karl Van Wyk, Umar Iqbal, Stan Birchfield, Jan Kautz, and Dieter Fox. DexYCB: A benchmark for capturing hand grasping of objects. In *CVPR*, 2021.
- [6] Yin Cui, Menglin Jia, Tsung-Yi Lin, Yang Song, and Serge Belongie. Class-balanced loss based on effective number of samples. In *CVPR*, 2019.
- [7] Dawson-Haggerty et al. Trimesh.
- [8] Jia Deng, Wei Dong, Richard Socher, Li-Jia Li, Kai Li, and Li Fei-Fei. ImageNet: A large-scale hierarchical image database. In *CVPR*, 2009.
- [9] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. BERT: Pre-training of deep bidirectional Transformers for language understanding. In *NAACL*, 2019.
- [10] Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, Jakob Uszkoreit, and Neil Houlsby. An image is worth 16x16 words: Transformers for image recognition at scale. In *ICLR*, 2021.
- [11] Zicong Fan, Adrian Spurr, Muhammed Kocabas, Siyu Tang, Michael J Black, and Otmar Hilliges. Learning to disambiguate strongly interacting hands via probabilistic per-pixel part segmentation. In *3DV*, 2021.
- [12] Zicong Fan, Omid Taheri, Dimitrios Tzionas, Muhammed Kocabas, Manuel Kaufmann, Michael J Black, and Otmar Hilliges. ARCTIC: A dataset for dexterous bimanual hand-object manipulation. In *CVPR*, 2023.
- [13] Yao Feng, Vasileios Choutas, Timo Bolkart, Dimitrios Tzionas, and Michael J Black. Collaborative regression of expressive bodies using moderation. In *3DV*, 2021.
- [14] Patrick Grady, Chengcheng Tang, Christopher D Twigg, Minh Vo, Samarth Brahmabhatt, and Charles C Kemp. ContactOpt: Optimizing contact to improve grasps. In *CVPR*, 2021.
- [15] Shreyas Hampali, Mahdi Rad, Markus Oberweger, and Vincent Lepetit. HOnnotate: A method for 3D annotation of hand and object poses. In *CVPR*, 2020.
- [16] Shreyas Hampali, Sayan Deb Sarkar, Mahdi Rad, and Vincent Lepetit. Keypoint Transformer: Solving joint identification in challenging hands and object interactions for accurate 3D pose estimation. In *CVPR*, 2022.
- [17] Mohamed Hassan, Vasileios Choutas, Dimitrios Tzionas, and Michael J Black. Resolving 3D human pose ambiguities with 3D scene constraints. In *ICCV*, 2019.
- [18] Mohamed Hassan, Partha Ghosh, Joachim Tesch, Dimitrios Tzionas, and Michael J Black. Populating 3D scenes by learning human-scene interaction. In *CVPR*, 2021.

- [19] Yana Hasson, Gul Varol, Dimitrios Tzionas, Igor Kalevatykh, Michael J Black, Ivan Laptev, and Cordelia Schmid. Learning joint reconstruction of hands and manipulated objects. In *CVPR*, 2019.
- [20] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *CVPR*, 2016.
- [21] Junxing Hu, Hongwen Zhang, Zerui Chen, Mengcheng Li, Yunlong Wang, Yebin Liu, and Zhenan Sun. Learning explicit contact for implicit reconstruction of hand-held objects from monocular images. In *AAAI*, 2024.
- [22] Chun-Hao P Huang, Hongwei Yi, Markus Höschle, Matvey Safroshkin, Tsvetelina Alexiadis, Senya Polikovsky, Daniel Scharstein, and Michael J Black. Capturing and inferring dense full-body human-scene contact. In *CVPR*, 2022.
- [23] Alexander Kirillov, Eric Mintun, Nikhila Ravi, Hanzi Mao, Chloe Rolland, Laura Gustafson, Tete Xiao, Spencer Whitehead, Alexander C. Berg, Wan-Yen Lo, Piotr Dollár, and Ross Girshick. Segment Anything. In *ICCV*, 2023.
- [24] Taein Kwon, Bugra Tekin, Jan Stühmer, Federica Bogo, and Marc Pollefeys. H2O: Two hands manipulating objects for first person interaction recognition. In *ICCV*, 2021.
- [25] Jaerin Lee, Daniel Sungho Jung, Kanggeon Lee, and Kyoung Mu Lee. SemanticDraw: Towards real-time interactive content creation from image diffusion models. In *CVPR*, 2025.
- [26] Zhaoqi Leng, Mingxing Tan, Chenxi Liu, Ekin Dogus Cubuk, Jay Shi, Shuyang Cheng, and Dragomir Anguelov. PolyLoss: A polynomial expansion perspective of classification loss functions. In *ICLR*, 2022.
- [27] Jiefeng Li, Siyuan Bian, Chao Xu, Zhicun Chen, Lixin Yang, and Cewu Lu. HybrIK-X: Hybrid analytical-neural inverse kinematics for whole-body mesh recovery. *TPAMI*, 2025.
- [28] Zhihao Li, Jianzhuang Liu, Zhensong Zhang, Songcen Xu, and Youliang Yan. CLIFF: Carrying location information in full frames into human pose and shape estimation. In *ECCV*, 2022.
- [29] Kevin Lin, Lijuan Wang, and Zicheng Liu. End-to-end human pose and mesh reconstruction with Transformers. In *CVPR*, 2021.
- [30] Kevin Lin, Lijuan Wang, and Zicheng Liu. Mesh graphormer. In *ICCV*, 2021.
- [31] Tsung-Yi Lin, Piotr Dollár, Ross Girshick, Kaiming He, Bharath Hariharan, and Serge Belongie. Feature pyramid networks for object detection. In *CVPR*, 2017.
- [32] Tsung-Yi Lin, Priya Goyal, Ross Girshick, Kaiming He, and Piotr Dollár. Focal loss for dense object detection. In *ICCV*, 2017.
- [33] Zhifeng Lin, Changxing Ding, Huan Yao, Zengsheng Kuang, and Shaoli Huang. Harmonious feature learning for interactive hand-object pose estimation. In *CVPR*, 2023.
- [34] Shaowei Liu, Yang Zhou, Jimei Yang, Saurabh Gupta, and Shenlong Wang. ContactGen: Generative contact modeling for grasp generation. In *ICCV*, 2023.
- [35] Yumeng Liu, Xiaoxiao Long, Zemin Yang, Yuan Liu, Marc Habermann, Christian Theobalt, Yuexin Ma, and Wenping Wang. EasyHOI: Unleashing the power of large models for reconstructing hand-object interactions in the wild. In *CVPR*, 2025.
- [36] Yunze Liu, Yun Liu, Che Jiang, Kangbo Lyu, Weikang Wan, Hao Shen, Boqiang Liang, Zhoujie Fu, He Wang, and Li Yi. HOI4D: A 4D egocentric dataset for category-level human-object interaction. In *CVPR*, 2022.
- [37] Matthew Loper, Naureen Mahmood, Javier Romero, Gerard Pons-Moll, and Michael J Black. SMPL: A skinned multi-person linear model. *ACM TOG*, 2015.
- [38] Ilya Loshchilov and Frank Hutter. Decoupled weight decay regularization. In *ICLR*, 2019.

- [39] Junzhe Lu, Jing Lin, Hongkun Dou, Ailing Zeng, Yue Deng, Xian Liu, Zhongang Cai, Lei Yang, Yulun Zhang, Haoqian Wang, and Ziwei Liu. DPoser-X: Diffusion model as robust 3D whole-body human pose prior. In *ICCV*, 2025.
- [40] Gyeongsik Moon, Shunsuke Saito, Weipeng Xu, Rohan Joshi, Julia Buffalini, Harley Bellan, Nicholas Rosen, Jesse Richardson, Mallorie Mize, Philippe de Bree, Tomas Simon, Bo Peng, Shubham Garg, Kevyn McPhail, and Takaaki Shiratori. A dataset of relighted 3D interacting hands. In *NeurIPS*, 2023.
- [41] Gyeongsik Moon, Shoou-I Yu, He Wen, Takaaki Shiratori, and Kyoung Mu Lee. InterHand2.6M: A dataset and baseline for 3D interacting hand pose estimation from a single RGB image. In *ECCV*, 2020.
- [42] Lea Muller, Ahmed AA Osman, Siyu Tang, Chun-Hao P Huang, and Michael J Black. On self-contact and human pose. In *CVPR*, 2021.
- [43] Hyeongjin Nam, Daniel Sungho Jung, Gyeongsik Moon, and Kyoung Mu Lee. Joint reconstruction of 3D human and object via contact-based refinement Transformer. In *CVPR*, 2024.
- [44] Hyeongjin Nam, Daniel Sungho Jung, Yeonguk Oh, and Kyoung Mu Lee. Cyclic test-time adaptation on monocular video for 3D human mesh reconstruction. In *ICCV*, 2023.
- [45] JoonKyu Park, Daniel Sungho Jung, Gyeongsik Moon, and Kyoung Mu Lee. Extract-and-adaptation network for 3D interacting hand mesh recovery. In *ICCV*, 2023.
- [46] JoonKyu Park, Yeonguk Oh, Gyeongsik Moon, Hongsuk Choi, and Kyoung Mu Lee. HandOccNet: Occlusion-robust 3D hand mesh estimation network. In *CVPR*, 2022.
- [47] Adam Paszke, Sam Gross, Francisco Massa, Adam Lerer, James Bradbury, Gregory Chanan, Trevor Killeen, Zeming Lin, Natalia Gimelshein, Luca Antiga, Alban Desmaison, Andreas Köpf, Edward Yang, Zach DeVito, Martin Raison, Alykhan Tejani, Sasank Chilamkurthy, Benoît Steiner, Lu Fang, Junjie Bai, and Soumith Chintala. PyTorch: An imperative style, high-performance deep learning library. In *NeurIPS*, 2019.
- [48] Georgios Pavlakos, Vasileios Choutas, Nima Ghorbani, Timo Bolkart, Ahmed AA Osman, Dimitrios Tzionas, and Michael J Black. Expressive body capture: 3D hands, face, and body from a single image. In *CVPR*, 2019.
- [49] Georgios Pavlakos, Dandan Shan, Ilija Radosavovic, Angjoo Kanazawa, David Fouhey, and Jitendra Malik. Reconstructing hands in 3D with Transformers. In *CVPR*, 2024.
- [50] Aditya Prakash, Benjamin Lundell, Dmitry Andreychuk, David Forsyth, Saurabh Gupta, and Harpreet Sawhney. How do i do that? Synthesizing 3D hand motion and contacts for everyday interactions. In *CVPR*, 2025.
- [51] Nikhila Ravi, Valentin Gabeur, Yuan-Ting Hu, Ronghang Hu, Chaitanya Ryali, Tengyu Ma, Haitham Khedr, Roman Rädle, Chloe Rolland, Laura Gustafson, Eric Mintun, Junting Pan, Kalyan Vasudev Alwala, Nicolas Carion, Chao-Yuan Wu, Ross Girshick, Piotr Dollár, and Christoph Feichtenhofer. SAM 2: Segment Anything in images and videos. In *ICLR*, 2025.
- [52] Tal Ridnik, Emanuel Ben-Baruch, Nadav Zamir, Asaf Noy, Itamar Friedman, Matan Protter, and Lihi Zelnik-Manor. Asymmetric loss for multi-label classification. In *ICCV*, 2021.
- [53] Robin Rombach, Andreas Blattmann, Dominik Lorenz, Patrick Esser, and Björn Ommer. High-resolution image synthesis with latent diffusion models. In *CVPR*, 2022.
- [54] Javier Romero, Dimitrios Tzionas, and Michael J Black. Embodied hands: Modeling and capturing hands and bodies together. *ACM TOG*, 2017.
- [55] Soshi Shimada, Vladislav Golyanik, Patrick Pérez, and Christian Theobalt. Decaf: Monocular deformation capture for face and hand interactions. *ACM TOG*, 2023.
- [56] Kihyuk Sohn, Honglak Lee, and Xinchen Yan. Learning structured output representation using deep conditional generative models. In *NeurIPS*, 2015.

- [57] Ke Sun, Bin Xiao, Dong Liu, and Jingdong Wang. Deep high-resolution representation learning for human pose estimation. In *CVPR*, 2019.
- [58] Shashank Tripathi, Agniv Chatterjee, Jean-Claude Passy, Hongwei Yi, Dimitrios Tzionas, and Michael J Black. DECO: Dense estimation of 3D human-scene contact in the wild. In *ICCV*, 2023.
- [59] Dimitrios Tzionas, Luca Ballan, Abhilash Srikantha, Pablo Aponte, Marc Pollefeys, and Juergen Gall. Capturing hands in action using discriminative salient points and physics simulation. *IJCV*, 2016.
- [60] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Lukasz Kaiser, and Illia Polosukhin. Attention is all you need. In *NeurIPS*, 2017.
- [61] Qingxuan Wu, Zhiyang Dou, Sirui Xu, Soshi Shimada, Chen Wang, Zhengming Yu, Yuan Liu, Cheng Lin, Zeyu Cao, Taku Komura, Vladislav Golyanik, Christian Theobalt, Wenping Wang, and Lingjie Liu. DICE: End-to-end deformation capture of hand-face interactions from a single image. In *ICLR*, 2025.
- [62] Xianghui Xie, Xi Wang, Nikos Athanasiou, Bharat Lal Bhatnagar, Chun-Hao P. Huang, Kaichun Mo, Hao Chen, Xia Jia, Zerui Zhang, Liangxian Cui, Xiao Lin, Bingqiao Qian, Jie Xiao, Wenfei Yang, Hyeongjin Nam, Daniel Sungho Jung, Kihoon Kim, Kyoung Mu Lee, Otmar Hilliges, and Gerard Pons-Moll. RHOBIN Challenge: Reconstruction of human object interaction. *arXiv preprint arXiv:2401.04143*, 2024.
- [63] Yufei Xu, Jing Zhang, Qiming Zhang, and Dacheng Tao. ViTPose: Simple Vision Transformer baselines for human pose estimation. In *NeurIPS*, 2022.
- [64] Lihe Yang, Bingyi Kang, Zilong Huang, Xiaogang Xu, Jiashi Feng, and Hengshuang Zhao. Depth Anything: Unleashing the power of large-scale unlabeled data. In *CVPR*, 2024.
- [65] Yifei Yin, Chen Guo, Manuel Kaufmann, Juan Jose Zarate, Jie Song, and Otmar Hilliges. Hi4D: 4D instance segmentation of close human interaction. In *CVPR*, 2023.
- [66] Hongwen Zhang, Yating Tian, Yuxiang Zhang, Mengcheng Li, Liang An, Zhenan Sun, and Yebin Liu. PyMAF-X: Towards well-aligned full-body model regression from monocular images. *TPAMI*, 2023.
- [67] Qitao Zhao, Ce Zheng, Mengyuan Liu, Pichao Wang, and Chen Chen. PoseFormerV2: Exploring frequency domain for efficient and robust 3D human pose estimation. In *CVPR*, 2023.

NeurIPS Paper Checklist

The checklist is designed to encourage best practices for responsible machine learning research, addressing issues of reproducibility, transparency, research ethics, and societal impact. Do not remove the checklist: **The papers not including the checklist will be desk rejected.** The checklist should follow the references and follow the (optional) supplemental material. The checklist does NOT count towards the page limit.

Please read the checklist guidelines carefully for information on how to answer these questions. For each question in the checklist:

- You should answer [Yes], [No], or [NA].
- [NA] means either that the question is Not Applicable for that particular paper or the relevant information is Not Available.
- Please provide a short (1–2 sentence) justification right after your answer (even for NA).

The checklist answers are an integral part of your paper submission. They are visible to the reviewers, area chairs, senior area chairs, and ethics reviewers. You will be asked to also include it (after eventual revisions) with the final version of your paper, and its final version will be published with the paper.

The reviewers of your paper will be asked to use the checklist as one of the factors in their evaluation. While "[Yes]" is generally preferable to "[No]", it is perfectly acceptable to answer "[No]" provided a proper justification is given (e.g., "error bars are not reported because it would be too computationally expensive" or "we were unable to find the license for the dataset we used"). In general, answering "[No]" or "[NA]" is not grounds for rejection. While the questions are phrased in a binary way, we acknowledge that the true answer is often more nuanced, so please just use your best judgment and write a justification to elaborate. All supporting evidence can appear either in the main paper or the supplemental material, provided in appendix. If you answer [Yes] to a question, in the justification please point to the section(s) where related material for the question can be found.

IMPORTANT, please:

- **Delete this instruction block, but keep the section heading “NeurIPS Paper Checklist”,**
- **Keep the checklist subsection headings, questions/answers and guidelines below.**
- **Do not modify the questions and only use the provided macros for your answers.**

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper’s contributions and scope?

Answer: [Yes]

Justification: We clearly state the core motivation to understand dense hand contact, outlines the challenges of class and spatial imbalance, introduces the proposed methods including balanced contact sampling and vertex-level class-balanced loss, and presents the contribution of a framework for learning dense hand contact estimation from imbalanced data. These claims are addressed in the introduction and supported by experimental results, accurately reflecting the paper’s contributions and scope.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: The limitations are discussed in the appendix, where we address the scope and boundaries of our method, including future directions.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[NA\]](#)

Justification: Our paper does not contain any theoretical result.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: The paper provides all necessary details to reproduce the main experiments, including data preprocessing, model architecture, training settings, and evaluation protocols.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: In the abstract of the paper, we provide URL to our publicly available code in an open repository.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.

- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [\[Yes\]](#)

Justification: The paper specifies dataset splits, model architecture, training schedules, optimizer settings, and other key hyperparameters necessary to understand and reproduce the results.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [\[No\]](#)

Justification: The results are obtained from standard randomly seeded runs without cherry-picking, offering a fair and representative evaluation despite not including error bars.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: The paper specifies the type of GPU used, batch size, number of training epochs, and key runtime settings needed to reproduce the experiments.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: The research adheres to the NeurIPS Code of Ethics, with no use of sensitive data or methods posing foreseeable ethical risks.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: The paper discusses societal impacts in the appendix, including both potential benefits and risks.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: The work does not involve releasing models or data with high risk of misuse, so specific safeguards are not applicable.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: All external datasets and code-bases used in the paper are properly cited.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: The paper does not introduce new assets such as datasets that require separate documentation.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.

- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: We mainly utilize publicly available datasets for our experiments.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: We mainly utilize publicly available datasets for our experiments.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. Declaration of LLM usage

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: LLM is not used in any critical component of the research process for this paper.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.

A Appendix

In this appendix, we provide additional technical details and experimental results that were omitted from the main manuscript due to space constraints. The contents are summarized below:

- A.1. Details of loss functions
- A.2. Details of contact initialization
- A.3. Details of dense hand contact labels
- A.4. More examples of class imbalance issue
- A.5. More examples of spatial imbalance issue
- A.6. More comparisons with state-of-the-art methods
- A.7. More ablation studies
- A.8. Quantitative results with different backbones
- A.9. Computational requirements
- A.10. More qualitative results
- A.11. Limitations and societal impacts

A.1 Details of loss functions

Smoothness loss. We penalize spatially isolated or fragmented contact predictions by leveraging the mesh topology of the hand. Let $p_v \in [0, 1]$ denote the predicted contact probability after the sigmoid for each vertex $v \in V$, and let $\mathbf{A} \in \{0, 1\}^{V \times V}$ be a sparse adjacency matrix representing vertex connectivity in the mesh. Here, $A_{vu} = 1$ indicates that vertex u is connected to vertex v , and 0 otherwise. For numerical stability, we also include self-connections so that each vertex is connected to itself. For each vertex $v \in V$, we compute the aggregated contact prediction over its neighbors as:

$$\hat{p}_v = \sum_{u \in V} A_{vu} \cdot p_u, \quad (9)$$

and similarly compute the aggregated non-contact prediction as:

$$\hat{q}_v = \sum_{u \in V} A_{vu} \cdot (1 - p_u), \quad (10)$$

where $\hat{p}_v$ and $\hat{q}_v$ represent the neighborhood-smoothed estimates for contact and non-contact probabilities, respectively. To maintain consistent scale across batches, we further normalize the aggregated contact and non-contact maps by their maximum values within each batch before computing the isolation score. To assess spatial inconsistency, we compute the discrepancy between each vertex’s prediction and its neighborhood average. Specifically, we define the isolation score as:

$$s_v = |p_v - \hat{p}_v| + |(1 - p_v) - \hat{q}_v|. \quad (11)$$

This term penalizes contact predictions that sharply differ from those of their neighbors, whether in the contact or non-contact region. To prevent the isolation score from disproportionately penalizing spatially large or confident regions, we normalize the total discrepancy by the sum of neighborhood weights:

$$n_v = \sum_{u \in V} A_{vu} = \hat{p}_v + \hat{q}_v, \quad (12)$$

which corresponds to the degree of vertex v in the mesh. This normalization accounts for how many vertices contribute to the isolation score at each location, ensuring that the loss measures average inconsistency rather than accumulating errors over larger neighborhoods. As a result, the model is encouraged to produce smooth contact maps without being biased against large spatially coherent contact regions. We then define the final smoothness regularization loss as:

$$\mathcal{L}_{\text{iso}} = \log \left(1 + \frac{\sum_{v \in V} s_v}{\sum_{v \in V} n_v + \epsilon} \right), \quad (13)$$

where ϵ is a small constant added for numerical stability. This loss encourages smooth transitions and penalizes sharp discontinuities in the prediction map. By suppressing highly localized contact signals and promoting coherent spatial clusters, it guides the model to favor fewer but larger contact regions, better reflecting real-world hand interaction patterns.

Regularization loss. To prevent the model from predicting arbitrarily skewed contact distributions, we introduce a global regularization loss that encourages predictions to stay close to the dataset-wide mean contact pattern. Let $p_v \in [0, 1]$ denote the predicted contact probability after the sigmoid for each vertex $v \in V$, and let $\bar{p}_v \in [0, 1]$ denote the dataset-wide average contact probability at vertex v , computed over all training samples. We then define the regularization loss as:

$$\mathcal{L}_{\text{reg}} = \frac{1}{|V|} \sum_{v \in V} |p_v - \bar{p}_v|, \quad (14)$$

which is an L1 loss between the predicted contact probabilities and the mean distribution. This encourages the model to avoid degenerate or overly confident predictions that deviate significantly from typical contact patterns, providing a soft global constraint during training.

A.2 Details of contact initialization

To stabilize the training of HACO on dense hand contact estimation, we propose a contact initialization technique that utilizes a learnable embedding to learn the most effective initial contact for dense hand contact estimation. In Table A1, our proposed contact initialization outperforms all other initialization methods, achieving the superior performance in precision, recall, and F1-score. Unlike other approaches that rely on fixed priors or simplistic assumptions, our method introduces a contact representation learned from input data, enabling more accurate and robust initialization during inference. This formulation captures the most useful initial prediction compared to static or handcrafted methods, leading to gains across all metrics. Such consistent improvement demonstrates that learning contact directly from large-scale datasets, rather than relying on pre-defined numbers or dataset averages, is crucial for precise contact estimation.

Table A1: **Comparison of various contact initialization on MOW [4] dataset.**

Methods	Precision $\uparrow$	Recall $\uparrow$	F1-Score $\uparrow$
No initialization	0.508	0.589	0.503
No contact	0.514	0.532	0.482
Full contact	0.511	0.551	0.492
Mean of DexYCB [5]	0.520	0.588	0.512
Mean of MOW [4]	0.511	<u>0.600</u>	<u>0.515</u>
Mean of IH26M [41]	<u>0.522</u>	0.567	0.505
Ours	0.525	0.607	0.522

A.3 Details of dense hand contact labels

Following the previous works on dense human-scene contact estimation [18, 22, 43], we employ distance-based thresholding to generate ground-truth dense hand contact labels. This is implemented using the Trimesh library [7]. Given a ground-truth hand mesh and the mesh of the interacting entity (i.e., an object for hand-object interaction, another hand for hand-hand interaction, a face for hand-face interaction, the environment for hand-scene interaction, or the body for hand-body interaction), we construct a proximity query from the interacting mesh and compute the closest surface point on it for each vertex of the hand mesh. Contact is then determined by thresholding the Euclidean distance between each hand vertex and its corresponding closest point. The resulting contact labels are binary values per vertex, indicating whether each vertex is in contact.

To set these thresholds, we manually inspected the ground-truth meshes and the resulting contact labels produced under fixed threshold values. We observed that the accuracy of 3D annotations varies significantly across datasets, making a single threshold unsuitable for all cases. A uniform threshold often led to incorrect labeling and semantically implausible contact regions. Based on our analysis, we set the contact distance threshold to 1 cm for the following datasets: ObMan [19], DexYCB [5], HO3D [15], H2O3D [16], ARCTIC [12], HOI4D [36], H2O [24], and PROX [17]. For the MOW [4]

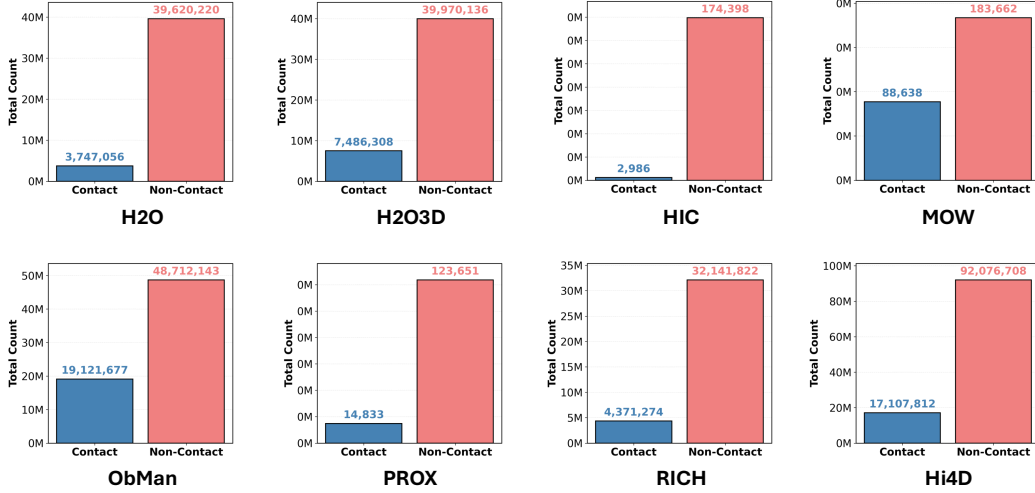


Figure A1: More examples of class imbalance between contact and non-contact on H2O [24], H2O3D [16], HIC [59], MOW [4], ObMan [19], PROX [17], RICH [22], Hi4D [65].

dataset, we used a threshold of 3.5 cm due to its lower mesh fidelity. For InterHand2.6M [41] and HIC [59], we adopted a smaller threshold of 0.5 cm to better capture fine-grained hand-hand interactions.

Although most datasets produced reasonable contact labels with a 1 cm threshold, MOW required a larger value of 3.5 cm due to coarse mesh annotations generated through optimization on in-the-wild images. For hand-hand interaction datasets, the 1 cm threshold, which is approximately the width of a finger, often resulted in inaccurate and semantically implausible contact labels. To address this, we adopted a smaller threshold of 0.5 cm. For datasets with existing annotated contact labels, such as RICH [22], Decaf [55], and Hi4D [65], we directly used the provided ground-truth annotations.

A.4 More examples of class imbalance issue

Figure A1 presents the class imbalance issue across datasets [24, 16, 59, 4, 19, 17, 22, 65] omitted from the main paper due to space constraints. In general, all datasets exhibit a severe class imbalance, with the majority of vertices corresponding to non-contact. Notably, some datasets such as H2O and HIC display highly skewed ratios of 10.57:1 and 58.41:1, respectively, between non-contact and contact vertices. It is important to note that contact and non-contact counts are computed at the vertex-level rather than per-hand instance. Since the ground-truth contact labels are provided per-hand, rather than per-vertex for each image, sampling must be performed at the hand-level. This makes it challenging to tackle class imbalance issue in per-hand manner as the class imbalance issue is present in per-vertex. Hence, it is essential to design a principled, hand-level proxy that reflects the underlying per-vertex imbalance. To tackle this challenge, we introduce balanced contact sampling (BCS), which constructs multiple sampling groups in hand-level but still fairly represent both contact and non-contact in per-vertex level. Please refer to the main paper for further details on the BCS.

A.5 More examples of spatial imbalance issue

We visualize the dataset-wide mean contact maps for all dense hand contact datasets used in training HACO in Figure A2 and Figure A3. Among hand-object interaction datasets, DexYCB [5], HO3D [15], H2O3D [16], ARCTIC [12], HOI4D [36], and H2O [24] exhibit significant spatial imbalance, with ground-truth contact concentrated predominantly on the fingertips. ObMan [19] and MOW [4] display more promising contact distributions, including high contact probability on the hypothenar region. However, ObMan is a synthetic dataset with domain gap towards real images, and MOW, while showing diverse palmar contact patterns, lacks dorsal contact and contains fewer than 1,000 samples, making it insufficient for large-scale training. Among hand-hand interaction datasets, HIC [59] shows severe spatial imbalance, with contact limited to small fingertip regions and

regions between fingers. In contrast, InterHand2.6M [41] offers a more favorable contact distribution, covering nearly all hand regions except for a small area on the dorsum. Hand-scene interaction datasets [17, 22] tend to involve large contact areas, primarily on the palmar side. However, interactions are often limited to flat surfaces such as ground or walls, resulting in low diversity of contact types despite broad coverage. The hand-face interaction dataset, Decaf [55], also exhibits poor spatial diversity due to its limited action types, such as poking or touching the chin. For the hand-body interaction dataset, Hi4D [65] includes extensive contact regions, but most interactions involve full palmar contact during actions such as hugging or patting. As visualized in the dataset-wide mean contact heatmaps, spatial imbalance is a prevalent issue across existing dense hand contact datasets. To address this, we propose the vertex-level class-balanced (VCB) loss, which mitigates spatial imbalance by reweighting the loss at each vertex based on its mean contact frequency across the dataset. Please refer to the main paper for further details on the VCB loss.

Table A2: **Comparison with SOTA methods of dense hand contact estimation on HIC [59], RICH [22], Hi4D [65] dataset.**

Test dataset	Methods	Precision $\uparrow$	Recall $\uparrow$	F1-Score $\uparrow$
HIC [59]	POSA [18]	0.000	0.000	0.000
	BSTRO [22]	0.000	0.000	0.000
	DECO [58]	<u>0.005</u>	<u>0.037</u>	<u>0.006</u>
	HACO (Ours)	0.216	0.409	0.263
RICH [22]	POSA [18]	0.143	0.175	0.125
	BSTRO [22]	<u>0.437</u>	<u>0.498</u>	<u>0.455</u>
	DECO [58]	0.324	0.351	0.303
	HACO (Ours)	0.741	0.899	0.781
Hi4D [65]	POSA [18]	0.128	0.109	0.111
	BSTRO [22]	<u>0.313</u>	<u>0.207</u>	<u>0.247</u>
	DECO [58]	0.197	0.162	0.146
	HACO (Ours)	0.555	0.636	0.565

A.6 More comparisons with state-of-the-art methods

Dense hand contact estimation models. In our main manuscript and Table A2, we provided comparison between our HACO and state-of-the-art (SOTA) methods [18, 22, 58] without any modification on model and its model weights. Nevertheless, the scale of training dataset is largely different between HACO and the SOTA methods. Also, all of the prior methods are dense human-scene contact estimation methods based on human body models such as SMPL [37] and SMPL-X [48]. To provide the comparison that excludes the effect of large-scale training and discrepancy of performance rooted by different human models (e.g., SMPL, SMPL-X, MANO), we further provide comparison with SOTA methods that are trained on the same 14 datasets like HACO and modified as hand version in Table A3. To modify the methods to hand version, we replaced all components of SMPL or SMPL-X modules within POSA [18], BSTRO [22], DECO [58] into MANO [54] modules and changed corresponding model configurations such as the number of vertices accordingly. For POSA, as the method requires estimated 3D mesh, we provide the estimated 3D hand mesh with HaMeR [49]. According to Table A3, HACO still consistently outperforms all hand-specific variants of the full-body methods. These results highlight the strength of HACO’s task-specific design based on balanced contact sampling and vertex-level class-balanced loss.

Dense hand contact estimation modules. While not designed as a standalone dense hand contact estimation model, there are several works that present dense hand contact estimation module as part of 3D reconstruction model [55, 21, 61, 35] or pseudo ground-truth generation of dense hand contact [50]. For contact modules in 3D reconstruction model, they can be divided into learning-based modules [55, 21, 61] and heuristic-driven techniques [35, 50]. These methods were excluded from our SOTA comparison in the main manuscript due to the lack of official code release [55, 21, 61], the task-specific scope of their training data [55, 61], and the heuristic nature of their modules [35, 50]. However, we include a broader evaluation in Table A4 to provide a more comprehensive perspective for future research and to facilitate extensive assessment of hand-specific methods. We first reproduce DefConNet from Decaf [55], Stage 1 of CHOI [21], and InteractionNet from DICE [61] based on the model specifications described in their manuscripts, as neither the code nor the pretrained weights are

Table A3: **Comparison with SOTA methods of dense hand contact estimation with hand version of POSA [18], BSTRO [22], DECO [58] on MOW [4], HIC [59], RICH [22], Hi4D [65] dataset.**

Test dataset	Methods	Precision $\uparrow$	Recall $\uparrow$	F1-Score $\uparrow$
MOW [4]	POSA [18] (hand ver. w/ HaMeR [49] pred. mesh)	0.504	0.400	0.408
	BSTRO [22] (hand ver.)	<u>0.518</u>	0.359	0.380
	DECO [58] (hand ver.)	0.495	0.385	0.395
	HACO (Ours)	0.525	0.607	0.522
HIC [59]	POSA [18] (hand ver. w/ HaMeR [49] pred. mesh)	0.060	0.225	0.086
	BSTRO [22] (hand ver.)	<u>0.070</u>	<u>0.293</u>	<u>0.107</u>
	DECO [58] (hand ver.)	0.067	0.210	0.094
	HACO (Ours)	0.216	0.409	0.263
RICH [22]	POSA [18] (hand ver. w/ HaMeR [49] pred. mesh)	0.538	0.355	0.398
	BSTRO [22] (hand ver.)	<u>0.581</u>	<u>0.413</u>	<u>0.421</u>
	DECO [58] (hand ver.)	0.519	0.321	0.367
	HACO (Ours)	0.741	0.899	0.781
Hi4D [65]	POSA [18] (hand ver. w/ HaMeR [49] pred. mesh)	0.537	0.257	0.315
	BSTRO [22] (hand ver.)	0.594	<u>0.375</u>	<u>0.420</u>
	DECO [58] (hand ver.)	0.553	0.244	0.304
	HACO (Ours)	<u>0.555</u>	0.636	0.565

publicly available. For the Data Engine in LatentAct [50] and EasyHOI [35], we use the officially released code to generate outputs for evaluation. For all learning-based modules [55, 21, 61], we train the reproduced modules based on 14 datasets like HACO but follow their training configurations such as loss function or optimizers when provided in their manuscripts. In Table A4, the results show that our HACO outperforms all previous dense hand contact estimation modules by a significant margin. This highlights the need for a standalone dense hand contact estimation model specifically trained for hand contact, and demonstrates the effectiveness of HACO’s task-specific design.

Table A4: **Comparison with SOTA methods of dense hand contact estimation with hand contact estimation modules (†: reproduced, re-trained with 14 datasets, code not available) on MOW [4] dataset.**

Methods	Precision $\uparrow$	Recall $\uparrow$	F1-Score $\uparrow$
†DefConNet (from Decaf [55])	0.414	<u>0.404</u>	0.368
†Stage 1 in CHOI [21]	<u>0.521</u>	0.387	0.407
†InteractionNet (from DICE [61])	0.459	0.392	<u>0.413</u>
Data Engine in LatentAct [50] (w/ SAM2 [51] & HaMeR [49])	0.345	0.244	0.211
EasyHOI [35]	0.480	0.228	0.282
HACO (Ours)	0.525	0.607	0.522

A.7 More ablation studies

Effectiveness of training dataset size. To assess the contribution of large-scale training to HACO’s performance, we conducted an ablation study comparing different training dataset sizes spanning 1, 3, and 14 datasets. The results demonstrate that performance improves notably as the training dataset size increases, emphasizing the importance of large-scale and diverse data for robust contact estimation. In particular, training on the full set of 14 datasets yields the best performance, while using only three representative datasets of DexYCB [5], ObMan [19], MOW [4] and even one dataset of MOW dataset still produces competitive results, indicating that HACO remains effective even with limited training dataset size.

A.8 Quantitative results with different backbones

We evaluate the performance of HACO with various backbone architectures, keeping the rest of the pipeline fixed. As shown in Table A6, our original model using ViT-H [10] pretrained on HaMeR [49] achieves the highest F1-score, demonstrating the strongest overall performance. ViT-based backbones consistently outperform convolutional alternatives, with ViT-B achieving the second highest recall

Table A5: Ablation for training dataset size on MOW [4] dataset.

Methods	Precision↑	Recall↑	F1-Score↑
HACO trained on 1 dataset	0.498	0.348	0.373
HACO trained on 3 datasets	0.463	<u>0.602</u>	<u>0.485</u>
HACO trained on 14 datasets	0.525	0.607	0.522

and ViT-S maintaining strong precision. ViT-L also performs competitively, achieving a strong F1-score. This highlights the effectiveness of Transformer [60]-based architectures for dense contact estimation, likely due to their global receptive fields and capacity to model long-range dependencies. To assess the benefit of pretraining on related tasks, we additionally train and evaluate a variant with FPN [31] backbone pretrained with HandOccNet [46], a 3D hand mesh reconstruction model. This configuration yields low F1-score compared to ImageNet-pretrained backbones, suggesting that reconstruction-specific pretraining alone is insufficient for accurate hand contact prediction. HRNet backbones, widely used in 3D human mesh reconstruction [29, 30, 28, 66, 27], achieve moderate performance in our setting, with F1-scores below 0.5 despite their high-resolution design. ResNet backbones [20] exhibit a wide performance range, with deeper variants such as ResNet-152 performing relatively well, while lighter variants such as ResNet-18 offer improved efficiency at the cost of reduced F1-score. Overall, these results indicate that the performance gains of our final HACO model arise from the combination of the strong representational capacity of the ViT backbone and task-specific knowledge transferred from 3D hand mesh reconstruction.

Table A6: Comparison of various backbone models on MOW [4] dataset.

Model	Backbone	Pretrained	Precision↑	Recall↑	F1-Score↑
HACO	ViT-H [10]	HaMeR [49]	0.525	0.607	0.522
HACO	FPN [31]	HandOccNet [46]	0.505	0.575	0.482
HACO	ViT-L [10]	ImageNet [8]	0.493	0.598	<u>0.491</u>
HACO	ViT-B [10]	ImageNet [8]	0.488	<u>0.604</u>	0.484
HACO	ViT-S [10]	ImageNet [8]	<u>0.510</u>	0.514	0.462
HACO	HRNet-W48 [57]	ImageNet [8]	0.509	0.563	0.485
HACO	HRNet-W32 [57]	ImageNet [8]	0.500	0.558	0.471
HACO	ResNet-152 [20]	ImageNet [8]	0.506	0.558	0.486
HACO	ResNet-101 [20]	ImageNet [8]	0.498	0.564	0.480
HACO	ResNet-50 [20]	ImageNet [8]	0.490	0.558	0.470
HACO	ResNet-34 [20]	ImageNet [8]	0.489	0.558	0.466
HACO	ResNet-18 [20]	ImageNet [8]	0.484	0.538	0.456

A.9 Computational requirements

Table A7 reports the computational requirements of HACO with various backbone configurations. The original HACO model, listed in the first row, shows the highest computational cost, particularly in training memory, which exceeds 26GB. This is primarily due to the memory overhead associated with storing gradients while fine-tuning the ViT-H [10] backbone pretrained on HaMeR [49]. To support broader adoption in downstream tasks such as 3D hand grasp generation [14] and 3D hand-object reconstruction [35], we provide additional HACO variants with reduced computational demand. We first evaluate the FPN [31] backbone from HandOccNet [46]. Although FPN is known to be computationally expensive due to its multi-scale feature hierarchy, it remains significantly more efficient than ViT-H. We also test lighter ViT variants such as ViT-B and ViT-S, which offer favorable trade-offs by maintaining moderate memory usage while achieving inference speeds exceeding 60 fps. We further evaluate HRNet [57] backbones, including HRNet-W48 and HRNet-W32. Despite their moderate memory consumption, both models exhibit low inference speeds, falling below 20 fps. Finally, we include ResNet-based [20] variants ranging from lightweight models such as ResNet-18 to deeper configurations such as ResNet-152. These models generally show low memory usage, fewer parameters, and high runtime efficiency. In particular, ResNet-18 achieves an inference speed exceeding 100 fps while maintaining a computational cost under 3 GFLOPs.

Table A7: **Computational requirements of various backbone configurations.**

Model	Backbone	Train Memory (MB)	Test Memory (MB)	Params. (M)	Speed (fps)	GFLOPs
HACO	ViT-H [63]	26,157	12,728	671.10	54.26	125.63
HACO	FPN [31]	6,887	3,002	59.56	66.71	9.94
HACO	ViT-L [63]	16,649	4,604	341.91	65.79	62.87
HACO	ViT-B [63]	7,767	2,940	122.84	96.80	18.54
HACO	ViT-S [63]	4,960	2,434	56.34	86.88	5.10
HACO	HRNet-W48 [57]	11,296	7,574	122.37	10.52	23.48
HACO	HRNet-W32 [57]	9,026	7,806	86.13	17.14	12.55
HACO	ResNet-152 [20]	9,252	5,096	103.05	29.84	15.93
HACO	ResNet-101 [20]	7,471	4,978	87.40	41.14	11.07
HACO	ResNet-50 [20]	5,919	4,826	68.41	73.25	6.21
HACO	ResNet-34 [20]	4,163	2,456	56.75	80.85	5.03
HACO	ResNet-18 [20]	3,785	2,372	46.64	111.23	2.61

A.10 More qualitative results

Additional qualitative comparisons of dense hand contact estimation between HACO and prior methods, including POSA [18], BSTRO [22], and DECO [58], are shown in Figures A4, A5, and A6, which correspond to the MOW [4], Hi4D [65], and both HIC [59] and RICH [22] datasets, respectively. Overall, HACO consistently outperforms existing approaches by a substantial margin. Unlike prior methods that frequently fail to predict contact even when the hand is clearly interacting, HACO reliably captures accurate dense hand contact. Among the baselines, DECO demonstrates stronger performance than POSA and BSTRO, likely as a result of being trained on in-the-wild datasets with more diverse contact distributions. However, it remains to be affected by class imbalance, resulting in a high rate of false negatives. BSTRO is more severely affected by this imbalance, frequently predicting no contact across the entire hand and producing overly coarse contact regions concentrated in the palmar area. This behavior likely stems from its reliance on human-scene interaction datasets during training. POSA also exhibits frequent false negatives and fails to detect contact even in visibly interacting regions, likely due to its dependence on a pose prior rather than direct contact supervision with visual input.

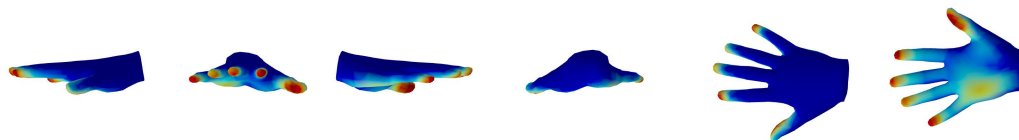
A.11 Limitations and societal impacts

Limitations. Our HACO supports dense hand contact estimation across a wide range of interaction types, including hand-object, hand-hand, hand-face, hand-scene, and hand-body interactions. However, hands frequently engage in self-contact under certain poses (e.g., contact between the thumb and index finger in an "okay" gesture) or during whole-body motions [42]. We intentionally exclude self-contact cases, as including them could negatively impact downstream applications. For example, ContactOpt [14] may incorrectly optimize grasps between the hand and an object if self-contact is misinterpreted as external contact. As a result, HACO currently does not support dense hand contact estimation in self-contact scenarios. Nevertheless, we consider modeling self-contact a promising direction for future research, particularly for applications in the metaverse (e.g., AR/VR) and action recognition. Furthermore, temporal context is essential to learn accurate dense hand contact estimation as contact is most noticeable when the movement of hand is obstructed by an object or scene. Integration of temporal-based methods [67, 44] that facilitate temporal information from input video can stabilize the performance of dense hand contact estimation. With the advent of diffusion-based methods [25, 39], designing a diffusion-based approach for dense hand contact estimation is another important research direction as it enables the learning of dense hand contact priors from the large-scale data. Additionally, our balanced contact sampling (BCS) utilizes a single reference point of dataset-wide dense hand contact mean. While this allows HACO to fairly represent diverse contact statistics, if distribution of the statistics is multi-modal, this strategy might not be effective. There can be a room for improvement on our BCS strategy if the multi-modality of dense hand contact datasets is proven and multi-modal aware methodology can be incorporated to BCS. Moreover, we evaluate HACO only on samples that contain at least one vertex in contact. However, this evaluation protocol may overlook the model's inability to handle false hand-level contact predictions, which are important to consider for real-world applications. The incorporation of additional evaluation protocols such as accuracy, specificity, and Matthews Correlation Coefficient may be

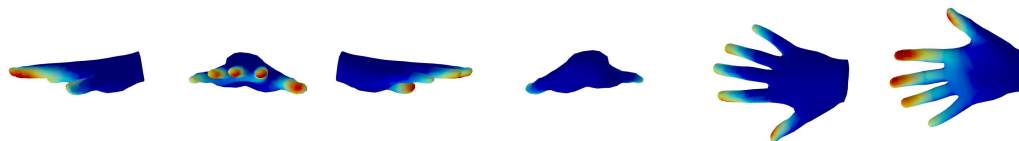
beneficial. Lastly, dense hand contact estimation demonstrates inferior performance on interacting hand scenarios relative to other hand interaction scenarios. The incorporation of 3D interacting hand mesh reconstruction method [11, 45] or the use of additional training dataset for hand-hand interaction [40] could enhance the performance.

Societal impacts. The proposed method has broad potential for applications involving hand interaction analysis, including AR/VR, robotics, and behavioral understanding. However, given the inherently human-centered nature of the task, there is a risk of misuse in areas such as mass surveillance or privacy infringement. We strongly discourage the deployment of this system in applications that may compromise human rights or personal privacy, and urge practitioners to consider ethical implications when applying HACO to downstream tasks.

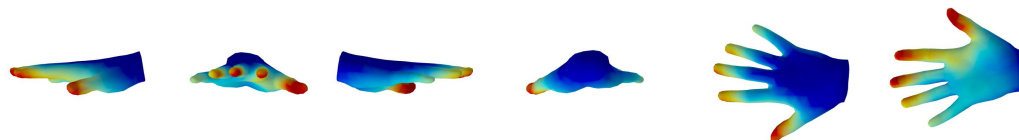
ObMan



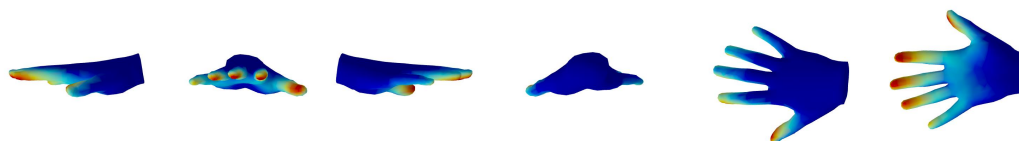
DexYCB



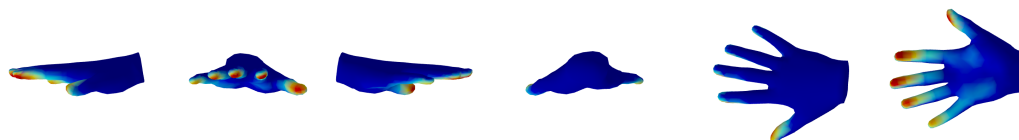
MOW



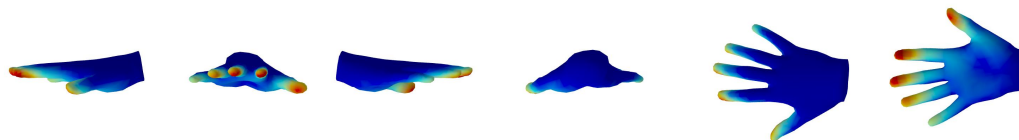
HO3D



H2O3D



ARCTIC



HOI4D

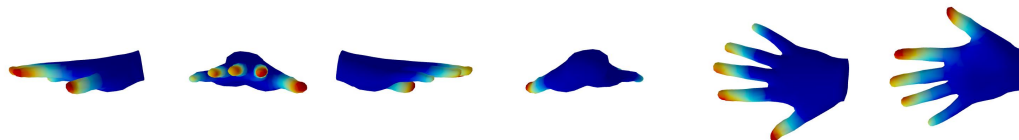
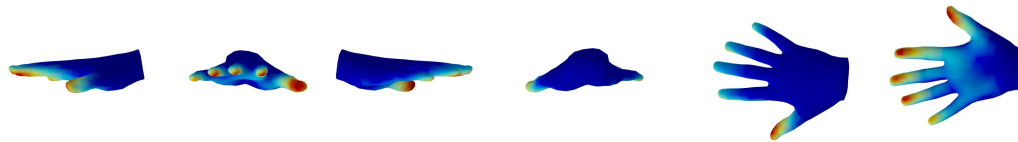
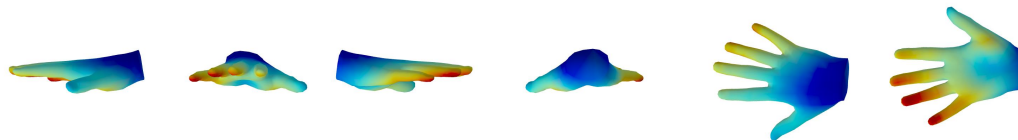


Figure A2: **Dataset-wise dense hand contact mean.** These heatmaps show average hand contact of ground-truths from ObMan [19], DexYCB [5], MOW [4], HO3D [15], H2O3D [16], ARCTIC [12], HOI4D [36] dataset.

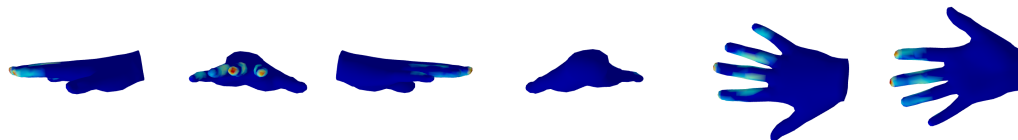
H2O



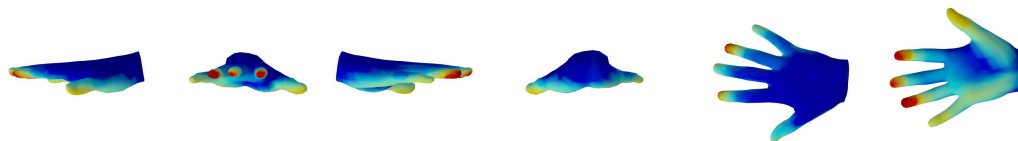
InterHand2.6M



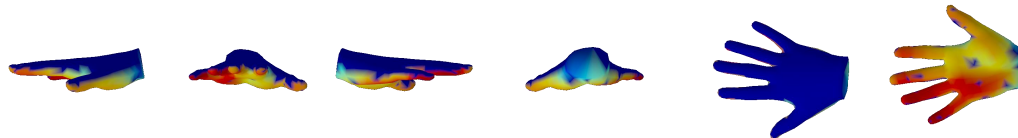
HIC



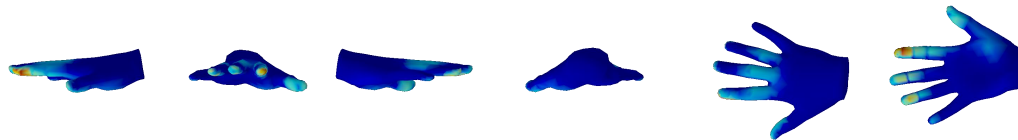
PROX



RICH



Decaf



Hi4D

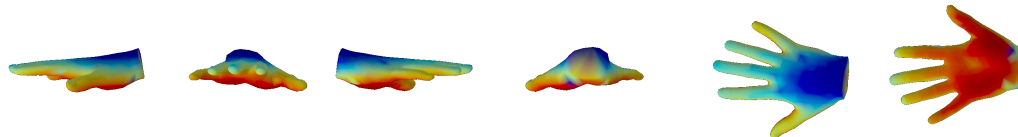


Figure A3: **Dataset-wise dense hand contact mean.** These heatmaps show average hand contact of ground-truths from H2O [24], InterHand2.6M [41], HIC [59], PROX [17], RICH [22], Decaf [55], Hi4D [65] dataset.

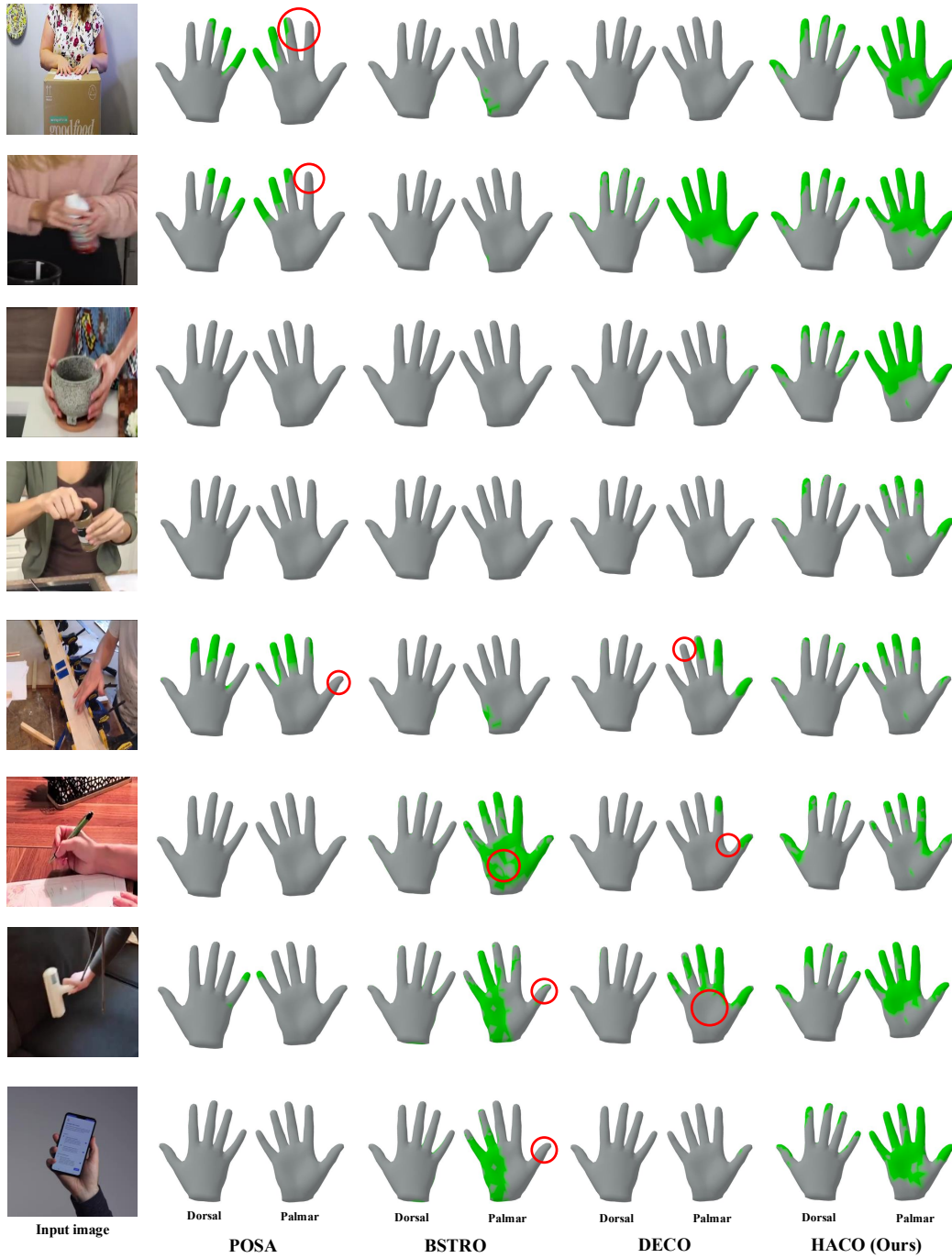


Figure A4: **Qualitative comparison of dense hand contact estimation with POSA [18], BSTRO [22], DECO [58] on MOW [4] dataset.** We highlight exemplar regions where HACO outperforms previous methods. Note that we only predict right hand contact.

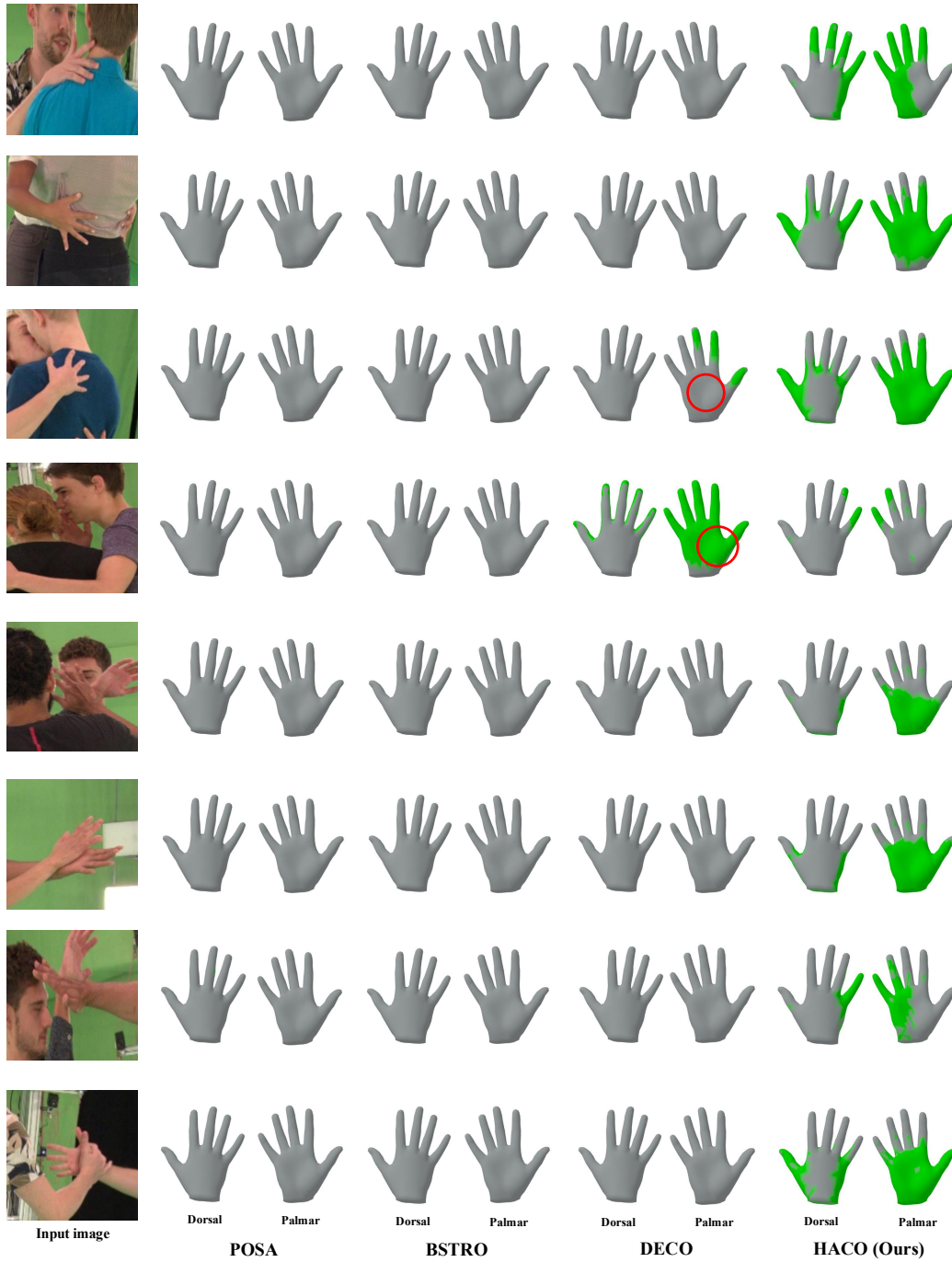


Figure A5: **Qualitative comparison of dense hand contact estimation with POSA [18], BSTRO [22], DECO [58] on Hi4D [65] dataset.** We highlight exemplar regions where HACO outperforms previous methods. Note that we only predict right hand contact.

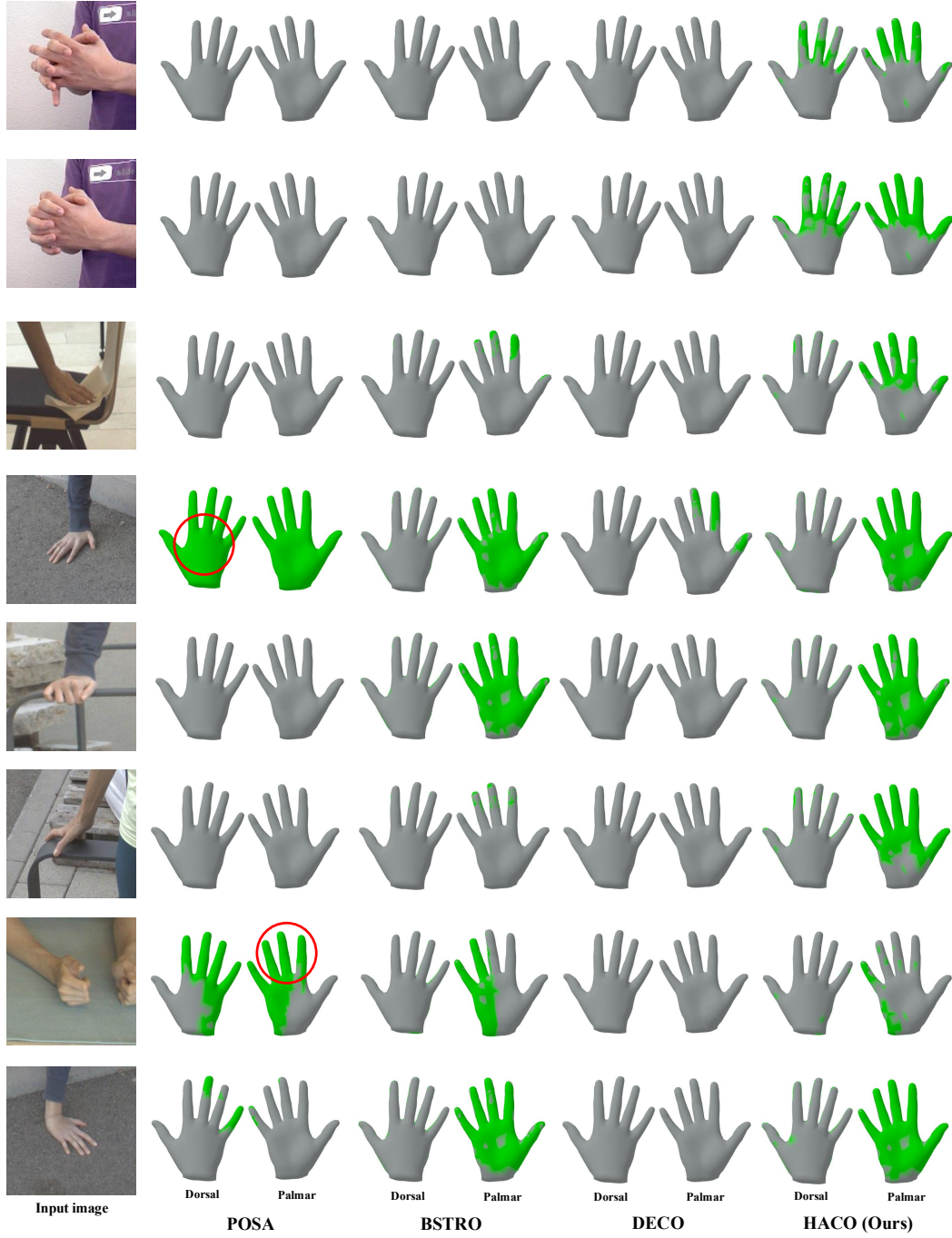


Figure A6: **Qualitative comparison of dense hand contact estimation with POSA [18], BSTRO [22], DECO [58] on HIC [59] and RICH [22] dataset.** We highlight exemplar regions where HACO outperforms previous methods. Note that we only predict right hand contact.