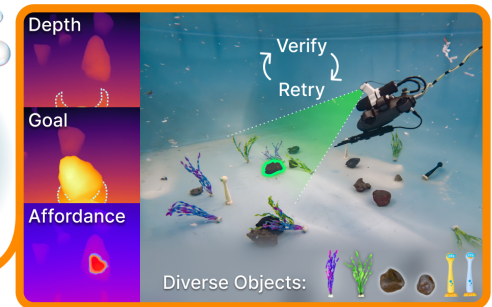


UMI-Underwater: Learning Underwater Manipulation without Underwater Teleoperation

Hao Li*, Long Yin Chung*, Jack Goler, Ryan Zhang, Xiaochi Xie, Huy Ha, Shuran Song, Mark Cutkosky

Task-centric Affordance on Land



Robot-Centric Self-Supervised Data in Water

Fig. 1. **UMI-Underwater.** We address data-collection and generalization bottlenecks in underwater manipulation by pairing autonomous, self-supervised data collection with a zero-shot, depth-based affordance predictor that transfers directly from land to water.

Abstract—Underwater robotic grasping is difficult due to degraded, highly variable imagery and the expense of collecting diverse underwater demonstrations. We introduce a system that (i) autonomously collects successful underwater grasp demonstrations via a self-supervised data collection pipeline and (ii) transfers grasp knowledge from on-land human demonstrations through a depth-based affordance representation that bridges the on-land-to-underwater domain gap and is robust to lighting and color shift. An affordance model trained on on-land handheld demonstrations is deployed underwater zero-shot via geometric alignment, and an affordance-conditioned diffusion policy is then trained on underwater demonstrations to generate control actions. In pool experiments, our approach improves grasping performance and robustness to background shifts, and enables generalization to objects seen only in on-land data, outperforming RGB-only baselines. Code, videos, and additional results are available at <https://umi-under-water.github.io>.

I. INTRODUCTION

Underwater robotic manipulation is key to enabling applications such as ecological sampling, debris removal, and infrastructure inspection, yet robust autonomy remains challenging. Compared to terrestrial settings, underwater robots operate under severe perceptual degradation and variability—including wavelength-dependent attenuation, scattering/backscatter, turbidity, and rapidly changing illumination and caustics—that can substantially shift appearance even within the same environment [1–4]. In addition, hydrodynamic disturbances and vehicle–object coupling complicate contact-rich interaction. These factors often make end-to-end visuomotor policies brittle under distribution shift, and they amplify the cost of collecting sufficient underwater training data.

A major practical bottleneck is the *human burden in data collection*. Most underwater manipulation systems remain

teleoperation-centric, but collecting diverse, high-quality underwater demonstrations is time-consuming and expensive. Recent work has begun to reduce this burden by leveraging self-supervised interaction to scale underwater data collection; by deploying autonomous heuristic controllers, systems can automatically acquire the successful demonstrations required to train robust visuomotor policies [5]. More broadly, self-supervised interaction has long been used to scale terrestrial robot data collection with automatically obtained success signals and repeated deployment [6–9].

In parallel, a growing line of work reduces demonstration burden by replacing robot-centric teleoperation with *portable, robot-agnostic human demonstration interfaces*. Universal Manipulation Interface (UMI) enables in-the-wild data collection with a handheld gripper and transfers learned visuomotor policies to robots [10]. Follow-up systems extend UMI-style data collection and transfer to new embodiments and deployment constraints, including mobile manipulation on legged platforms, aerial manipulation under challenging dynamics, and more scalable/hardware-independent variants with larger datasets and richer sensing [11–17].

This paper introduces a representation-centric method for robust underwater manipulation built around *affordance*, while explicitly targeting human burden reduction in two complementary ways. (1) We introduce a self-supervised underwater data collection pipeline that autonomously collects successful grasp demonstrations to eliminate reliance on teleoperation. The collection system bootstraps grasp attempts with a heuristic controller, executes recovery behaviors to increase data efficiency, and filters episodes using an automatic success signal, yielding scalable underwater demonstrations. (2) We develop a UMI-derived handheld interface (**UMI-Aquatic**) to scalably collect diverse on-land demonstrations to train a

*Equal contribution. All authors are with Stanford University. Emails: {li2053, lychung, jgoler, rrzhang, xxie15, huyha, shuran, cutkosky}@stanford.edu. Correspondence to: li2053@stanford.edu.

depth-conditioned affordance model for grasp guidance. This allows us to (i) achieve *zero-shot* land-to-water grasp knowledge transfer and (ii) bridge the land-to-water perception gap. Finally, we train an affordance- and depth-conditioned diffusion policy [18] on action trajectories from our autonomous underwater collection system. This visuomotor policy captures the perception robustness of depth and the diverse grasp knowledge of affordance to improve generalization to novel manipulation targets and unseen environments.

We evaluate our approach on two regimes: (i) novel-object generalization and (ii) visual generalization to unseen backgrounds and lighting conditions. Across these settings, we find that depth-based affordances provide an effective perception interface for bridging land-to-water shift, and that autonomous data collection substantially reduces the human burden required to train underwater manipulation policies.

Contributions:

- **Self-supervised underwater data collection:** a practical pipeline in a controlled pool setting that autonomously gathers successful underwater grasp demonstrations using recovery behaviors and automatic success filtering.
- **Depth-based affordance transfer from land to water:** a goal-conditioned, depth-based affordance predictor trained with on-land UMI-Aquatic demonstrations, enabling *zero-shot* land-to-water transfer without mixed training or fine-tuning.
- **Robustness and transfer evaluation:** experiments on in-distribution pool grasping, background shifts, and novel-object generalization using objects seen only in the on-land dataset.

II. RELATED WORK

A. Underwater Manipulation and Learning

Autonomous underwater manipulation has long been studied in the context of intervention missions (e.g., inspection, maintenance, and recovery), where the dominant challenges arise from hydrodynamic disturbances, limited visibility, and the complexity of controlling coupled vehicle–manipulator systems. Early work and benchmark systems emphasized reliable mechatronic integration and robust control for free-floating intervention, including SAUVIM-style intervention platforms [19]. Subsequent efforts developed increasingly capable autonomy stacks for underwater vehicle–manipulator systems (UVMS), including the TRIDENT framework and its descendants [20], as well as large-scale projects such as DexROV that advanced semi-autonomous intervention and remote supervision [21]. A complementary line of work focuses on principled control architectures, e.g., task-priority control for underwater intervention, to systematically handle competing objectives such as station keeping, collision avoidance, and end-effector motion [22].

A complementary paradigm is human–robot haptic collaboration: Ocean One and the deep-sea OceanOneK use an “avatar” telepresence approach that enables dexterous bimanual manipulation under supervisory control, with haptic feedback as needed [23–25].

Despite steady progress, many practical underwater manipulation systems remain teleoperation-centric, in part because collecting high-quality underwater demonstrations is expensive and time-consuming. Learning from demonstration has therefore often been used to assist or partially automate underwater intervention, for example by learning task models that help disambiguate operator intent and improve resilience to communication constraints [26]. More recently, AquaBot showed that end-to-end visuomotor policies can be trained from demonstrations and improved through iterative self-learning beyond initial teleoperation performance [5].

Building off this trajectory, our work addresses two complementary pillars of underwater manipulation: we replace expert teleoperation with an autonomous, self-supervised collection pipeline to scale data, while simultaneously exploring how transferable affordance representations can improve generalization across the land-to-water domain gap.

B. Reducing human burden in demonstration collection

A central challenge in learning-based manipulation is collecting diverse, high-quality demonstrations. In underwater settings this burden is amplified by limited visibility, vehicle–object coupling, and the operational overhead of teleoperation, motivating approaches that reduce (i) the amount of direct human-in-the-loop control and (ii) the effort required to collect diverse demonstrations.

a) Autonomous and self-supervised on-robot data collection: A long line of work demonstrates that robots can scale interaction data by autonomously executing trials and using automatically obtained success signals for learning, substantially reducing manual labeling and teleoperation time [6–9]. Recent systems revisit this theme with practical autonomous data-collection pipelines that repeatedly deploy policies, accumulate experience, and update models under real-world constraints [27–29]. Our underwater data collection pipeline follows this spirit: it autonomously generates grasp trials with recovery behaviors and retains successful episodes for behavior cloning, reducing reliance on underwater teleoperation.

b) Portable handheld interfaces and cross-embodiment transfer (UMI family): Complementary to autonomy-on-robot, handheld demonstration interfaces aim to make data collection cheap, portable, and robot-agnostic by capturing task-relevant trajectories with consistent observation/action semantics. Universal Manipulation Interface (UMI) enables in-the-wild demonstration collection using handheld grippers and transfers the learned visuomotor policies to robots without requiring robots at collection sites [10]. Subsequent UMI-style systems extend this paradigm along multiple axes: transferring UMI policies to new embodiments with additional control constraints (e.g., legged mobile manipulation and aerial manipulation) [11, 12], redesigning the interface for easier deployment and larger-scale datasets [13, 14], augmenting the interface with contact-rich sensing (e.g., force/torque and tactile modalities) [15, 30], and broadening the interface concept to dexterous hands via wearable or vision-based adaptations [17]. Our **UMI-Aquatic** interface is inspired by this line,

but targets a distinct challenge: leveraging cheap on-land demonstrations to improve underwater manipulation, where teleoperation is especially time-consuming. We use a depth-based affordance representation and geometric alignment to enable zero-shot land-to-water transfer without mixed training or fine-tuning.

C. Affordance Representations for Manipulation

Affordances provide task-relevant, spatially localized cues for interaction. Early deep affordance approaches predict pixel-wise affordance labels or masks jointly with object detection [31]. Dense action-value or utility maps over image pixels have also been used to guide grasping and interaction [8]. Recent work studies affordances that generalize to novel objects, including few-shot affordance learning for unseen articulated objects [32]. In parallel, robust manipulation can benefit from strong intermediate visual representations such as monocular depth [33, 34]. Recent underwater work also studies uncertainty-aware perceptual grasp selection, e.g., PUGS [35]. Our contribution is to use an affordance heatmap as a modular perception interface that can be trained with on-land data, and then used to guide an underwater diffusion policy.

III. APPROACH

A. Experimental Setup

Fig. 2 illustrates the overall system. The ROV manipulation setup is similar to that in prior autonomous underwater manipulation (e.g., [5]), with updates enabled by recent improvements to the QYSEA platform SDK.

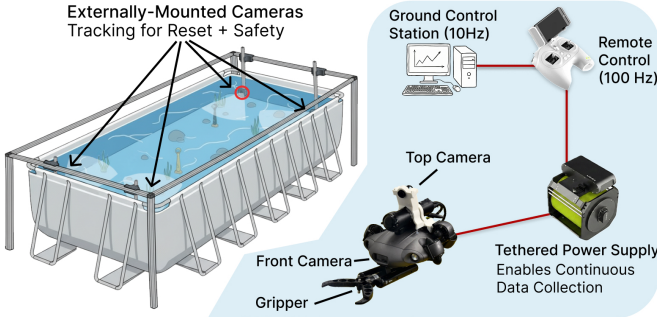


Fig. 2. **ROV Manipulation Setup for Self-supervised Underwater Data Collection.** The setup includes a swimming pool, with objects scattered across the workspace for repeated grasp attempts. Four fixed external cameras are mounted at the pool corners to provide real-time 3D localization used for safety functions (but not provided as inputs to the learned policy). The tethered ROV operates in this environment to execute the staged grasping routine.

a) *Platform:* The platform is a compact tethered underwater ROV equipped with a forward-facing RGB camera, an additional wide-FOV RGB camera mounted at the top, and a 1-DoF gripper (open/close). The vehicle is operated through a tethered control box at 100 Hz.

b) *Cameras:* The top camera in Fig. 2 is the externally mounted waterproof camera. We use two RGB views on the vehicle. The onboard front camera provides a forward-facing stream transmitted through the tether with below-100 ms latency. We additionally mount an external waterproof wide-FOV camera to provide a complementary viewpoint with richer scene context during approach and grasp.

c) *State estimation and safety:* The environment is equipped with four fixed external cameras placed at the pool corners for real-time 3D localization of the vehicle. Fusing the externally estimated 3D position with onboard IMU yields a 6-DoF pose estimate in a global frame. This position estimate is *not* provided to the learned policy; it is reserved for safety (collision avoidance), episode reset, and operator monitoring.

d) *Control constraints:* During both data collection and policy execution, vehicle pitch is fixed to a constant downward angle (10°) to stabilize the camera viewpoint and simplify the grasping geometry. This constraint reduces control complexity and improves repeatability while remaining representative of close-range grasping behaviors in this setting.

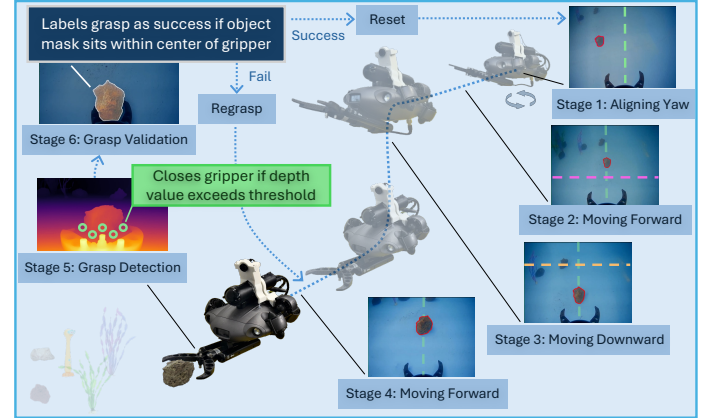


Fig. 3. **Autonomous Data Collection Pipeline.** Our heuristic controller autonomously collects grasping episodes by using a segmentation model to select a target, servoing the object centroid to stage-specific pixel setpoints using PD control, closing the gripper when a depth threshold is met, and labeling success via drag validation.

B. Data Collection

A key challenge in underwater learning is acquiring large quantities of robot interaction data without relying on extensive teleoperation. We therefore design a self-supervised data collection pipeline that uses a heuristic visuo-servoing PD controller to autonomously generate grasp trajectories. Each episode proceeds through a sequence of stages (Fig. 3).

a) *Reset and initialization:* At the start of each episode, the vehicle is placed at a shallow depth with a random initial pose that provides a wide field of view over the workspace. Using the external tracking system, we initialize the vehicle state and enforce the fixed pitch constraint.

b) *Object selection and tracking:* We run real-time segmentation with (SAM2-tiny [36]) to obtain per-frame object masks. From the visible objects, the robot autonomously selects the target whose centroid is closest to the image center and track its mask across frames. We compute a target pixel as the centroid of the selected mask.

c) *Stage 1: Yaw alignment:* We yaw the vehicle such that the image centerline coincides with the target centroid. Specifically, we define the horizontal pixel error between the centroid and the image centerline and use a PD controller to command yaw until the error falls below a threshold. It remains active throughout the grasp sequence to maintain target centering during approach and closure.

d) *Stage 2: Forward approach:* After initial yaw alignment, we command forward motion to reduce the apparent distance to the object. Concretely, we define a reference horizontal line in the lower portion of the image and continue moving forward until the centroid reaches this reference line (i.e., until the vertical pixel error meets a threshold). This provides a simple but effective proxy for range without requiring a calibrated stereo system.

e) *Stage 3: Depth adjustment:* We then adjust vehicle depth to keep the target within a favorable grasping region. We define a second reference horizontal line in the upper portion of the image and command vertical motion until the centroid lies within the target band.

f) *Stages 4 & 5: Close-range approach and grasp:* Once the vehicle is sufficiently close, we use a monocular depth estimator (Depth Anything V2 [37]) on the onboard RGB stream to further regulate approach distance. When the estimated distance indicates that the object is within the gripper workspace, the gripper closes to attempt a grasp.

g) *Stage 6: Drag verification:* To verify grasp success without external instrumentation, we execute a short retreat motion (dragging the object) for 3 s. If the object remains in the gripper without slipping during this interval, we label the episode as a success; otherwise as a failure. We use drag verification rather than lift-based verification because lift-based verification is about $3\times$ slower on our platform due to its limited payload margin. Since this step is used only for success labeling, it improves collection efficiency without affecting policy learning.

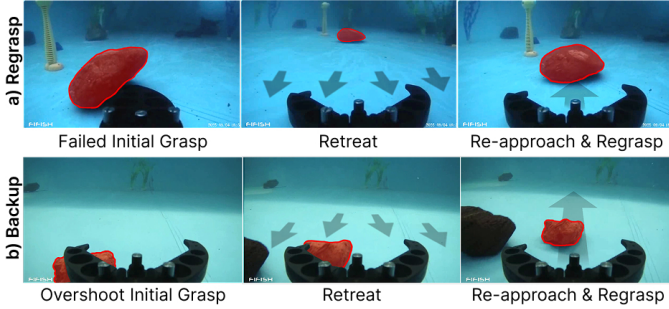


Fig. 4. **Autonomous Recovery Strategies** include (a) **regrasp** after failed grasps and (b) **backup** when the robot overshoots. These strategies improve the success rate of demonstration collection as well as improve policy robustness by demonstrating recovery behavior.

h) *Failure recovery mechanisms:* We incorporate two recovery behaviors that substantially improve autonomous collection (Fig. 4): (1) **Regrasp**. After each failed grasp attempt, the vehicle backs up, applies a small lateral offset (randomly left/right), reopens the gripper and retries the approach. This increases the probability of completing a successful grasp after poor initial contact. (2) **Overshoot**. Due to hydrodynamic effects, the vehicle can overshoot, moving the object toward the image boundary or out of view. We detect imminent loss-of-view when the target centroid approaches the image margin; in that case we command a larger retreat to re-acquire the object before restarting alignment.

i) *On-land demonstrations with UMI-Aquatic:* To collect diverse grasping demonstrations on land, we introduce

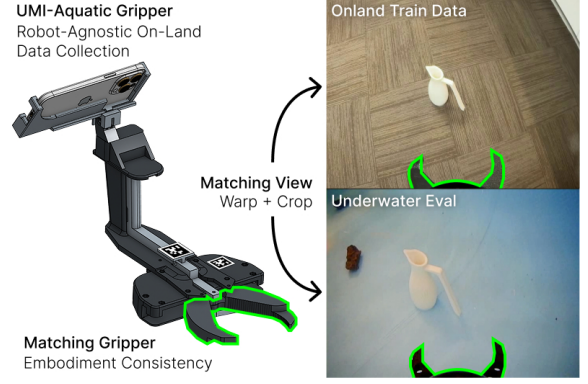


Fig. 5. **UMI-Aquatic on-land demonstration setup.** Our handheld gripper with an iPhone camera system and AprilTags enables portable data collection and reliable gripper-state tracking for automatic demonstration labeling. Cropping and geometric warping via reprojection align the iPhone view to match the underwater robot camera.

UMI-Aquatic, a handheld gripper interface derived from the Universal Manipulation Interface (UMI) [10]. UMI-Aquatic integrates (i) an iPhone camera rigidly mounted to match the underwater camera viewpoint and (ii) AprilTags mounted on the gripper to enable reliable gripper-state tracking during handheld operation (Fig. 5). Compared to the original UMI design, our modifications are tailored to underwater transfer: the camera mounting is designed to simplify cross-camera alignment, and the tag-based tracking provides a robust estimate of gripper openness/closure timing used for automatic label generation.

j) *Scale and human effort:* Our self-supervised pool pipeline collected 536 grasp demonstrations, with each episode lasting ~ 45 s and between-episode reset taking ~ 60 s, for a total of ~ 15 h of autonomous runtime for the episodes plus resets. From this set, we used 233 successful underwater grasps to train the diffusion policy.

Objects were re-scattered roughly every 5 h, and manual intervention was only required when the vehicle became stuck or the tether tangled. Separately, UMI-Aquatic enabled fast on-land collection of 800 handheld demonstrations across 6 object types and 8 backgrounds (each ~ 10 s, ~ 2.2 h of active demonstrating)

C. Goal-Conditioned Data Preprocessing and Training

We improve robustness and generalization by separating *where to grasp* from *how to control*. We learn (i) a goal-conditioned affordance predictor that outputs a dense grasp heatmap and (ii) a visuomotor policy that conditions on this heatmap (and depth) to produce control commands. To reduce sensitivity to lighting and color differences between land and underwater scenes, our affordance predictor operates on depth rather than RGB; in our setting, RGB-input affordance models trained on on-land UMI-Aquatic data transfer poorly to underwater imagery due to severe appearance shift.

a) *Training tuples and goal-conditioning:* From each episode we construct training tuples $\{(D_t, D_{\text{goal}}, \hat{H}_t)\}_{t=1}^T$, where D_t is the depth observation at time t , D_{goal} is a goal depth image, and $\hat{H}_t$ is a sparse per-pixel affordance target. Goal-conditioning disambiguates which object to grasp in multi-object scenes.

Deployment goal source. At deployment, the goal observation is provided **from on-land UMI-Aquatic data**: at test time we specify the target with a single on-land goal RGB image of the object inside the handheld gripper, convert it to a depth map using the same monocular depth pipeline, and use it as D_{goal} (after the same crop/warp preprocessing as in Sec. C. d). Thus, the affordance model performs **cross-domain goal matching** between the current underwater observation and an on-land goal, and we do **not** require capturing an underwater goal image at deployment.

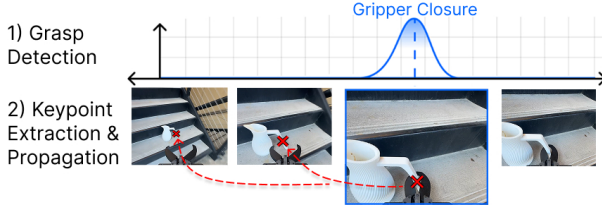


Fig. 6. **Automatic Affordance Supervision.** AprilTag tracking identifies gripper closure timing using a sliding window, which seeds point tracking to backtrack grasp contact keypoints on the object for affordance supervision.

b) Automatic per-frame supervision: For on-land data collected with UMI-Aquatic as shown in Fig. 5, we use an AprilTag-based gripper tracking pipeline (Fig. 6) to estimate gripper openness/width and to determine the grasp/contact timing in the camera frame. We detect the closure time t^* using a sliding-window change detector over the width signal (a sharp drop followed by a stable plateau). Then, we backtrack the gripper-object contact pixel from the first frame where the gripper is closed, t^* , to earlier frames using point tracking (CoTracker3 [38]), yielding a tracked contact location (u_t, v_t) for each $t \leq t^*$, producing compatible targets $\hat{H}_t$.

c) Land-to-water pretraining via diverse on-land data (zero-shot transfer): To improve generalization, we train an affordance model on a diverse on-land dataset. An iPhone-based capture setup is portable and easy to deploy, enabling rapid collection of varied interactions without specialized site preparation (e.g., no mapping run). Because the iPhone camera differs from the underwater camera in intrinsics, distortion, and field-of-view, we reduce this geometric domain gap by calibrating both camera models and warping on-land frames into the underwater camera geometry. Using depth as the affordance input further reduces the perception gap caused by lighting and color differences. Importantly, there is *no mixed training and no fine-tuning*: online affordance deployment uses on-land data only while offline affordance used for training the policy uses underwater data only.

d) Affordance dataset preparation and cross-camera warping: A practical challenge for land-to-water transfer is the mismatch between the iPhone camera used for on-land collection and the waterproof wide-view camera used underwater. We therefore pre-warp each on-land frame into the underwater camera image geometry using a calibrated, plane-at-depth remapping.

Plane-at-depth warp (iPhone $\rightarrow$ underwater camera). Let camera 1 denote the iPhone (source) and camera 2 denote

the underwater camera (target), with pinhole intrinsics K_1, K_2 and distortion parameters $\mathbf{d}_1, \mathbf{d}_2$. For each target pixel $\tilde{\mathbf{p}}_2 = (u_2, v_2)$ in camera 2, we undistort to normalized coordinates (x_2, y_2) and form a ray $\mathbf{r}_2 = [x_2, y_2, 1]^\top$. We intersect this ray with a fronto-parallel plane at depth Z in camera 2 coordinates:

$$\mathbf{X}_2 = Z \mathbf{r}_2. \quad (1)$$

We then transform to camera 1 coordinates and project into the iPhone image:

$$\mathbf{X}_1 = R \mathbf{X}_2 + t, \quad (2)$$

$$\tilde{\mathbf{p}}_1 = \pi(K_1, \mathbf{d}_1, \mathbf{X}_1), \quad (3)$$

where $\pi(\cdot)$ denotes perspective projection with lens distortion. Since we mount the iPhone camera to match the underwater camera orientation, we use $R = I$ (translation-only alignment). Finally, we perform backward sampling:

$$I_{1 \rightarrow 2}(\tilde{\mathbf{p}}_2) = I_1(\tilde{\mathbf{p}}_1), \quad (4)$$

implemented as a dense remapping table applied with bilinear resampling for images.

e) Affordance model architecture and training: We train a goal-conditioned affordance predictor from depth sequences. For each training sample, we take the current depth image D_t and a goal depth image D_{goal} , normalize depth to $[0, 1]$ using dataset-wide anchored min/max statistics (with clamping), resize to 112×112 , and concatenate (D_t, D_{goal}) into a two-channel input. Targets are per-pixel keypoint maps representing ground truth grasp locations, resized with nearest-neighbor interpolation. Our model is a U-Net (depth 4, base channels 64) with convolution-BN-ReLU encoder blocks and max-pooling, a bottleneck with increased channel capacity, and a mirrored decoder with transposed-convolution upsampling and skip connections. A final 1×1 convolution outputs per-pixel logits at the input resolution. Training uses a weighted BCE-with-logits objective to address foreground/background imbalance, AdamW optimization with warmup+cosine scheduling, gradient clipping, and mixed precision, with depth-specific augmentations (geometric transforms plus scale/noise/dropout/blur). We use episode-wise train/validation splits to avoid temporal leakage.

D. Affordance Conditioned Visuomotor Policy

a) Policy inputs: The diffusion policy receives a multi-modal input with: (i) predicted affordance heatmap, (ii) monocular depth map predicted from RGB, and (iii) low-dimensional robot status signals (compass, pitch, vehicle depth), as shown in Fig. 7. During policy training, we generate affordance heatmaps by training an affordance predictor on the same self-supervised underwater training dataset and applying it offline to all demo frames. At deployment, however, the online affordance heatmap is produced by the UMI-Aquatic affordance model trained only on on-land demonstrations and transferred underwater zero-shot. We do not use external motion-capture pose for policy input; it is reserved for safety and reset.

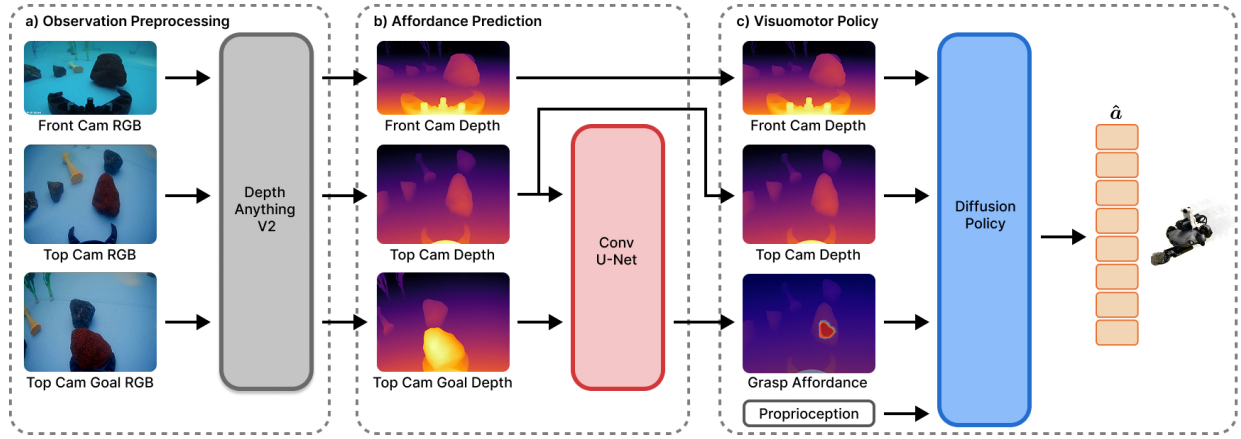


Fig. 7. **Training-time model architecture for the diffusion policy.** (a) During diffusion-policy training, current and goal observations are taken from the underwater robot top/front cameras and converted to depth maps using Depth Anything V2. (b) A training-time affordance predictor takes the top-camera current depth and top-camera goal depth observations and predicts a dense grasp affordance heatmap for target localization. This predictor is trained/applied offline on the self-supervised underwater training set to generate affordance inputs for policy training. (c) The visuomotor policy conditions on the predicted affordance map, depth observations, and robot proprioception to generate a short-horizon action trajectory. At deployment, the policy receives the same input format, but the affordance heatmap is produced by a separate UMI-Aquatic affordance model trained only on on-land demonstrations, using an on-land goal image as the target and transferring underwater zero-shot.

b) Diffusion policy: We adopt a diffusion policy that models a horizon of future actions and generates commands via iterative denoising. Observations are encoded by a pretrained vision backbone (ViT-B/16 CLIP [39]) into a global latent condition, and a 1D UNet denoiser predicts a sequence of actions. We predict a 16-step action horizon with a 6-dimensional action vector corresponding to yaw, forward/backward, up/down, left/right, open, close, executed in a receding-horizon manner. Given our perception and inference pipeline introduces approximately 200 ms of latency, this receding-horizon formulation provides smooth control between replanning steps while mitigating compounding errors.

c) Training and compute: We train the diffusion policy by behavior cloning on successful demonstrations using the standard diffusion objective (MSE on predicted noise for noised action trajectories). Training is performed on a single NVIDIA RTX 5090 GPU and AMD Ryzen 9800X3D CPU with batch size 56 for 150 epochs.

d) Inference and deployment: At test time, we run the affordance predictor, monocular depth estimation, and diffusion policy on the same device as training. The full perception-control stack runs online and publishes actions at 10 Hz, with policy inference updates at 1 Hz.

IV. EXPERIMENTS

A. Evaluation Setup

a) Environment: All experiments are conducted in an above-ground swimming pool, approximately 2×4 m in area and 1 m deep. Unless otherwise stated, objects are placed within a fixed workspace region on the pool floor and the robot executes a pick-and-drag grasping episode of fixed duration.

b) Task: Our primary task is *pick-and-drag grasping*: the robot must grasp a target object and maintain grasp while retreating for a fixed duration.

c) Multi-object evaluation and goal specification: In each trial, we place three objects in front of the robot (for both ID and OOD settings). All three objects are visible in

the robot’s initial camera view. We specify the desired target using a single on-land goal image from UMI-Aquatic: the operator selects an on-land goal RGB image of the intended target object, which is then converted to a goal depth map D_{goal} using the same monocular depth pipeline and preprocessing described in Sec. III-C. This goal is consumed by the *affordance model* to produce a goal-aligned heatmap used by affordance-conditioned policies. RGB-only baselines do not receive the goal image.

The main challenge is therefore not only executing a grasp, but also selecting the correct target among distractors under appearance variation and cross-domain goal specification (on-land → underwater).

d) Object sets and evaluation settings: We define a seen underwater object set used to collect pool demonstrations and train policies (rocks, seagrass, toy ducks), and a novel on-land UMI object set used only for evaluating novel-object transfer (vase with handle, can, power drill). We evaluate three settings:

- 1) **In-distribution (ID):** seen underwater objects under the *training* pool background distribution.
- 2) **OOD visual generalization (background shift):** seen underwater objects under *unseen* pool background wall-papers.
- 3) **OOD novel-object generalization (UMI objects):** on-land UMI objects evaluated in the pool under the *training* pool background distribution. This evaluation is zero-shot. There is no policy fine-tuning at this evaluation.

e) Metric: We report success rate (%; higher is better): an episode is successful if the robot grasps the specified target object and it does not slip during a 3 s retreat/drag verification (Section III-B). Grasping a non-target object is counted as a failure. For each method and condition, we run 20 trials.

B. Methods Compared

We compare diffusion-policy controllers trained on successful pool demonstrations, with different perceptual inputs and

affordance pretraining sources. Let **RGB** denote raw RGB observations, **Depth** denote monocular depth predicted from RGB, and **Aff** denote the learned goal-conditioned affordance heatmap. All diffusion policies are trained on the same pool demonstration set (seen underwater objects under the training pool background), and differ only in their observation modalities and/or the affordance model used. We focus on these within-platform comparisons because prior underwater manipulation systems typically differ in hardware, sensing, task structure, and teleoperation assumptions, making direct reproduction a poor way to isolate the contribution of the perception interface in our setting.

a) *Goal inputs and fairness:* For methods that use **Aff**, the goal image is provided *only* to the affordance model to produce the heatmap; the diffusion policy itself is conditioned on the predicted heatmap (plus depth and proprioception), not directly on the goal image. The **DP+RGB** and **DP+Depth** baselines are **goal-agnostic** and represent the common setting of training end-to-end underwater policies from in-water demonstrations without cross-domain goal-conditioning. Since these baselines receive no explicit goal specification in multi-object scenes, we define their implicit target as the object closest to the image center in the initial view, matching the center-biased target selection used during our self-supervised data collection pipeline (Section III-B). Accordingly, the improvement of **DP+Aff+Depth** over the goal-agnostic baselines should be interpreted as reflecting both explicit goal-conditioning and the robustness of the affordance/depth interface. To isolate the effect of raw RGB while holding goal-conditioning fixed, the controlled comparison in our experiments is **DP+Aff+Depth** versus **DP+RGB+Aff+Depth**, which share the same goal-conditioned affordance input and differ only by the addition of RGB. To isolate the contribution of depth under background shift, we additionally compare against **DP+Depth**.

b) *Important note on affordance supervision:* Although it is possible to train an affordance model using only pool-collected self-supervised data, in our setting this pool-only dataset is highly biased: our underwater data collection follows a multi-stage heuristic controller that repeatedly approaches objects with similar trajectories and viewpoints, and the dataset is relatively small and dominated by a narrow subset of objects (e.g., rocks). As a result, pool-only affordance training produces less robust affordance predictions. Our method uses UMI-Aquatic to collect on-land demonstrations of the pool object categories (rocks, seagrass, toy ducks), trains the affordance model on this dataset, and then deploys it underwater in a zero-shot manner.

c) *Compared methods:*

- **DP + Aff + Depth (ours) (UMI-Aquatic-pretrained Aff):** diffusion policy conditioned on the affordance heatmap and monocular depth (and proprioception), trained on pool demonstrations from the seen-object set; the affordance model is trained on on-land UMI-Aquatic demonstrations and transferred underwater zero-shot (no mixed training, no fine-tuning).
- **DP + RGB (goal-agnostic; seen-only):** diffusion pol-

TABLE I
ID PERFORMANCE (SEEN OBJECTS + TRAINING POOL BACKGROUND).

Method	Success (%) $\uparrow$
DP + Aff + Depth (ours)	17/20 (85%)
DP + RGB	13/20 (65%)

icy conditioned on RGB (and proprioception), trained on pool demonstrations from the seen-object set (rocks/seagrass/duck).

- **DP + Depth (goal-agnostic depth-only ablation):** diffusion policy conditioned on monocular depth (and proprioception), trained on the same pool demonstrations. We evaluate this baseline under unseen backgrounds to isolate how much robustness comes from depth alone, without affordance-based goal localization.
- **DP + RGB + Aff + Depth (RGB-addition ablation; UMI-Aquatic-pretrained Aff):** diffusion policy conditioned on RGB, affordance heatmap, and depth (and proprioception), trained on pool demonstrations; the affordance model is the same UMI-Aquatic-pretrained model as above.

d) *Modality ablations under background shift:* We include **DP + RGB + Aff + Depth** as an ablation to test whether providing *raw RGB* in addition to the robust (Aff+Depth) interface helps or hurts. We additionally include **DP + Depth** to measure how much appearance robustness comes from replacing RGB with monocular depth alone. Because the brittleness of RGB is most pronounced under large appearance shift, we primarily report these ablations under the unseen-background setting (Table II).

e) *Train/test protocol:* All diffusion policies are trained on pool demonstrations collected with the seen underwater object set under the training pool background distribution. We evaluate: (i) in-distribution (ID) performance on seen objects under the training background, (ii) OOD visual generalization on seen objects under unseen background wallpapers, and (iii) OOD novel-object generalization on novel objects placed in the pool (zero-shot; no policy fine-tuning at test time).

C. Results on the In-Distribution Task and Environment

We first evaluate in-distribution performance on the seen underwater object set under the training pool background, establishing a baseline before any distribution shift.

As shown in Table I, in the ID setting, **DP + Aff + Depth** outperforms the RGB-only baseline. The primary failure mode of **DP + RGB** in multi-object scenes is target confusion: it sometimes approaches and grasps a distractor object instead of the intended one. This suggests that a substantial part of the ID improvement comes from explicit goal-conditioned localization provided by the affordance heatmap.

D. OOD Visual Generalization Under Unseen Backgrounds

Table II shows that under background shift, RGB-containing policies fail catastrophically. The training pool background is visually simple and dominated by blue tones, whereas

TABLE II
OOD GENERALIZATION (SEEN OBJECTS + UNSEEN BACKGROUNDS).

Method	Success (%) $\uparrow$
DP + Aff + Depth (ours)	16/20 (80%)
DP + Depth	12/20 (60%)
DP + RGB	0/20 (0%)
DP + RGB + Aff + Depth	0/20 (0%)

the wallpaper backgrounds introduce strong texture and color statistics (e.g., wood-grain patterns) that lie outside the training distribution (see supp. material). As a result, **DP + RGB** and **DP + RGB + Aff + Depth** both achieve 0% success. In contrast, **DP + Depth** reaches 60%, showing that monocular depth alone already removes much of the appearance dependence. **DP + Aff + Depth** further improves to 80%, indicating that explicit goal-conditioned affordance localization provides additional benefit beyond depth by helping the policy select the correct target among distractors. Notably, **DP + RGB + Aff + Depth** uses the *same* affordance model and the *same* depth input as **DP + Aff + Depth**; its collapse to 0% indicates that adding RGB can cripple robustness under strong appearance shift even when other robust signals are available.

E. Novel-Object Generalization via On-Land Pretraining

We evaluate novel-object generalization by testing in the pool on objects that appear only in the on-land UMI dataset (vase with handle, can, power drill), which are disjoint from the underwater seen-object set used to train diffusion policies (rocks/seagrass/duck). Our method uses an affordance model trained on on-land UMI grasping data and transfers it underwater zero-shot, while the DP+RGB baseline is trained only on seen underwater objects.

As reported in Table III, on novel objects, **DP + Aff + Depth** substantially outperforms **DP + RGB**. Interestingly, the RGB-only policy still achieves non-trivial success (50%), which we attribute to two factors: (i) the background remains unchanged from training, so many low-level visual statistics match the training distribution; and (ii) the novel objects share some coarse geometric properties with the seen set, allowing occasional grasps despite lacking explicit goal-aligned localization. However, the affordance-conditioned policy is markedly more robust, suggesting that on-land UMI-pretrained affordances provide a stronger inductive bias for selecting graspable regions on previously unseen objects.

F. Failure mode: overshoot-induced target switch.

While the affordance heatmap substantially reduces target confusion in multi-object scenes, we observed a recurring failure mode when the robot *overshoots* during an approach. In these cases, the camera viewpoint can shift such that the target object moves partially out of view or becomes severely misaligned with the original goal specification. The affordance model may then place high confidence on a different salient region that remains visible (often a distractor object), which in turn guides the policy toward the wrong object. This failure can also *invalidate our recovery strategy*: because recovery

TABLE III
OOD NOVEL-OBJECT GENERALIZATION (UNSEEN OBJECTS IN POOL + TRAINING POOL BACKGROUND), ZERO-SHOT.

Method	Success (%) $\uparrow$
DP + Aff + Depth (ours)	15/20 (75%)
DP + RGB	10/20 (50%)

assumes the affordance remains anchored to the intended target, an overshoot-induced target switch can cause recovery to re-approach the distractor rather than returning to the original goal. We expect this issue can be mitigated by improved low-level control (reducing overshoot), and/or by adding temporal consistency to target localization (e.g., goal tracking or history-conditioned affordance prediction).

V. DISCUSSION AND LIMITATIONS

Our experiments are limited to pick-and-drag grasping in a controlled pool, where objects rest on a relatively flat support surface. This setting simplifies geometry estimation, approach planning, and support reasoning, and may reduce the data required to learn effective behavior. In more realistic underwater environments, uneven terrain, clutter, and greater scene variability may make localization and grasp execution substantially more difficult.

Our method also relies heavily on predicted depth. Although depth helps bridge the appearance gap between on-land and underwater RGB images, performance may degrade when monocular depth is unreliable or when geometry alone misses grasp-relevant cues. This trade-off may be especially limiting for objects with similar shapes but different appearances, or when success depends on appearance-specific features.

Future work should evaluate beyond flat-floor pool settings and include more varied underwater geometries, backgrounds, and support conditions in both training and testing. Combining depth with appearance cues, for example through RGB-depth fusion or uncertainty-aware conditioning, may further improve robustness. In addition, replacing our simple PD controller with a dynamics-aware controller such as MPC may reduce overshoot, improve tracking, and lessen the viewpoint shifts that contribute to target-switch failures.

VI. CONCLUSION

We presented a practical pipeline for underwater manipulation that combines self-supervised underwater data collection with a goal-conditioned affordance model enabled by UMI-Aquatic. The central idea is to separate *goal localization* from *control*: the affordance model predicts a goal-aligned spatial heatmap, and a policy conditioned on affordance and depth executes pick-and-drag grasping. This design improves in-distribution performance by reducing target confusion in multi-object scenes and increases robustness to large appearance shifts, such as unseen pool wallpapers. By transferring an affordance model trained on on-land data, our method also generalizes to novel underwater objects without policy fine-tuning, highlighting the value of on-land pretraining when paired with appropriate conditioning.

ACKNOWLEDGMENTS

This work was supported in part by the NSF Award #2143601, #2037101, and #2132519, and Samsung. We thank Tian-Ao Ren, Juhyun Jung, Stanley Wang, and Tianyu Tu for their thoughtful discussions and help with the experiments. We also give special thanks to Stanford Hopkins Marine Station for supporting our in-the-wild ocean deployment tests. The views and conclusions contained herein are those of the authors and should not be interpreted as necessarily representing the official policies, either expressed or implied, of the sponsors.

REFERENCES

- [1] Derya Akkaynak and Tali Treibitz. A revised underwater image formation model. In *2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 6723–6732, 2018.
- [2] Derya Akkaynak and Tali Treibitz. Sea-Thru: A method for removing water from underwater images. In *2019 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 1682–1691, 2019.
- [3] Chongyi Li, Chunle Guo, Wenqi Ren, Runmin Cong, Junhui Hou, Sam Kwong, and Dacheng Tao. An underwater image enhancement benchmark dataset and beyond. *IEEE Transactions on Image Processing*, 29:4376–4389, 2020.
- [4] Saeed Anwar and Chongyi Li. Diving deeper into underwater image enhancement: A survey. *Signal Processing: Image Communication*, 89:115978, 2020.
- [5] Ruoshi Liu, Huy Ha, Mengxue Hou, Shuran Song, and Carl Vondrick. Self-improving autonomous underwater manipulation. In *2025 IEEE International Conference on Robotics and Automation (ICRA)*, pages 16915–16922, 2025.
- [6] Lerrel Pinto and Abhinav Gupta. Supersizing self-supervision: Learning to grasp from 50k tries and 700 robot hours. In *2016 IEEE International Conference on Robotics and Automation (ICRA)*, pages 3406–3413, 2016.
- [7] Sergey Levine, Peter Pastor, Alex Krizhevsky, Julian Ibarz, and Deirdre Quillen. Learning hand-eye coordination for robotic grasping with deep learning and large-scale data collection. *The International Journal of Robotics Research*, 37(4–5):421–436, 2018.
- [8] Andy Zeng, Shuran Song, Stefan Welker, Johnny Lee, Alberto Rodriguez, and Thomas Funkhouser. Learning synergies between pushing and grasping with self-supervised deep reinforcement learning. In *2018 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, 2018.
- [9] Andy Zeng, Shuran Song, Johnny Lee, Alberto Rodriguez, and Thomas Funkhouser. Tossingbot: Learning to throw arbitrary objects with residual physics. *IEEE Transactions on Robotics*, 36(4):1307–1319, 2020.
- [10] Cheng Chi, Zhenjia Xu, Chuer Pan, Eric Cousineau, Benjamin Burchfiel, Siyuan Feng, Russ Tedrake, and Shuran Song. Universal manipulation interface: In-the-wild robot teaching without in-the-wild robots. *arXiv preprint arXiv:2402.10329*, 2024.
- [11] Huy Ha, Yihuai Gao, Zipeng Fu, Jie Tan, and Shuran Song. Umi-on-legs: Making manipulation policies mobile with manipulation-centric whole-body controllers. In *Proceedings of The 8th Conference on Robot Learning*, volume 270 of *Proceedings of Machine Learning Research*, pages 5254–5270. PMLR, 2025.
- [12] Harsh Gupta, Xiaofeng Guo, Huy Ha, Chuer Pan, Muqing Cao, Dongjae Lee, Sebastian Scherer, Shuran Song, and Guanya Shi. Umi-on-air: Embodiment-aware guidance for embodiment-agnostic visuomotor policies. *arXiv preprint arXiv:2510.02614*, 2025.
- [13] Zhaxizhuoma, Kehui Liu, Chuyue Guan, Zhongjie Jia, Ziniu Wu, Xin Liu, Tianyu Wang, Shuai Liang, Penggan CHEN, Pingrui Zhang, Haoming Song, Delin Qu, Dong Wang, Zhigang Wang, Nieqing Cao, Yan Ding, Bin Zhao, and Xuelong Li. Fastumi: A scalable and hardware-independent universal manipulation interface with dataset. In Joseph Lim, Shuran Song, and Hae-Won Park, editors, *Proceedings of The 9th Conference on Robot Learning*, volume 305 of *Proceedings of Machine Learning Research*, pages 3069–3093. PMLR, 27–30 Sep 2025.
- [14] Kehui Liu, Zhongjie Jia, Yang Li, Penggan Chen, Song Liu, Xin Liu, Pingrui Zhang, Haoming Song, Xinyi Ye, Nieqing Cao, et al. Fastumi-100k: Advancing data-driven robotic manipulation with a large-scale umi-style dataset. *arXiv preprint arXiv:2510.08022*, 2025.
- [15] Hojung Choi, Yifan Hou, Chuer Pan, Seongheon Hong, Austin Patel, Xiaomeng Xu, Mark R Cutkosky, and Shuran Song. In-the-wild compliant manipulation with umi-ft. *arXiv preprint arXiv:2601.09988*, 2026.
- [16] Omar Rayyan, John Abanes, Mahmoud Hafez, Anthony Tzes, and Fares Abu-Dakka. Mv-umi: A scalable multi-view interface for cross-embodiment learning. *arXiv preprint arXiv:2509.18757*, 2025.
- [17] Mengda Xu, Han Zhang, Yifan Hou, Zhenjia Xu, Linxi Fan, Manuela Veloso, and Shuran Song. Dexumi: Using human hand as the universal manipulation interface for dexterous manipulation. *arXiv preprint arXiv:2505.21864*, 2025.
- [18] Cheng Chi, Siyuan Feng, Yilun Du, Zhenjia Xu, Eric Cousineau, Benjamin CM Burchfiel, and Shuran Song. Diffusion policy: Visuomotor policy learning via action diffusion. In *Proceedings of Robotics: Science and Systems (RSS)*, Daegu, Republic of Korea, July 2023.
- [19] Giacomo Marani, Song K. Choi, and Junku Yuh. Underwater autonomous manipulation for intervention missions AUVs. *Ocean Engineering*, 36(1):15–23, 2009. Autonomous Underwater Vehicles.
- [20] Pedro J. Sanz, Pere Ridao, Gabriel Oliver, Claudio Melchiorri, Giuseppe Casalino, Carlos Silvestre, Yvan Petillot, and Alessio Turetta. Trident: A framework for autonomous underwater intervention missions with

- dexterous manipulation capabilities. *IFAC Proceedings Volumes*, 43(16):187–192, 2010. 7th IFAC Symposium on Intelligent Autonomous Vehicles.
- [21] Paolo Di Lillo, Enrico Simetti, Daniela De Palma, Elisabetta Cataldi, Giovanni Indiveri, Gianluca Antonelli, and Giuseppe Casalino. Advanced ROV autonomy for efficient remote control in the DexROV project. *Marine Technology Society Journal*, 50:67–80, 07 2016.
- [22] Enrico Simetti, Giuseppe Casalino, Francesco Wanderlingh, and Michele Aicardi. Task priority control of underwater intervention systems: Theory and applications. *Ocean Engineering*, 164:40–54, 09 2018.
- [23] Oussama Khatib, Xiyang Yeh, Gerald Brantner, Brian Soe, Boyeon Kim, Shameek Ganguly, Hannah Stuart, Shiquan Wang, Mark Cutkosky, Aaron Edsinger, Phillip Mullins, Mitchell Barham, Christian R. Voolstra, Khaled Nabil Salama, Michel L’Hour, and Vincent Creuze. Ocean One: A Robotic Avatar for Oceanic Discovery. *IEEE Robotics and Automation Magazine*, 23(4):20–29, December 2016.
- [24] Gerald Brantner and Oussama Khatib. Controlling Ocean One: Human-robot collaboration for deep-sea manipulation. *J. Field Robotics*, 38(1):28–51, 2021.
- [25] Hannah Stuart, Shiquan Wang, Oussama Khatib, and Mark R Cutkosky. The Ocean One hands: An adaptive design for robust marine manipulation. *The International Journal of Robotics Research*, 36(2):150–166, 2017.
- [26] Ioannis Havoutis and Sylvain Calinon. Learning from demonstration for semi-autonomous teleoperation. *Autonomous Robots*, 43(3):713–726, 2019.
- [27] Konstantinos Bousmalis, Giulia Vezzani, Dushyant Rao, Coline Devin, Alex X Lee, Maria Bauzá, Todor Davchev, Yuxiang Zhou, Agrim Gupta, Akhil Raju, et al. Robocat: A self-improving generalist agent for robotic manipulation. *arXiv preprint arXiv:2306.11706*, 2023.
- [28] Georgios Papagiannis and Edward Johns. Miles: Making imitation learning easy with self-supervision. In *Conference on Robot Learning*, pages 810–829. PMLR, 2025.
- [29] Suvir Mirchandani, Suneel Belkhale, Joey Hejna, Evelyn Choi, Md Sazzad Islam, and Dorsa Sadigh. So you think you can scale up autonomous robot data collection? In *Conference on Robot Learning*, pages 2791–2810. PMLR, 2025.
- [30] Tailai Cheng, Kejia Chen, Lingyun Chen, Liding Zhang, Yue Zhang, Yao Ling, Mahdi Hamad, Zhenshan Bing, Fan Wu, Karan Sharma, et al. Tacumi: A multi-modal universal manipulation interface for contact-rich tasks. *arXiv preprint arXiv:2601.14550*, 2026.
- [31] Thanh-Toan Do, Anh Nguyen, Ian Reid, Darwin G. Caldwell, and Nikos G. Tsagarakis. Affordancenet: An end-to-end deep learning approach for object affordance detection. In *2018 IEEE International Conference on Robotics and Automation (ICRA)*, pages 1–5, 2018.
- [32] Chuanruo Ning, Ruihai Wu, Haoran Lu, Kaichun Mo, and Hao Dong. Where2explore: Few-shot affordance learning for unseen novel categories of articulated objects. In *Advances in Neural Information Processing Systems*, pages 4585–4596, 2023.
- [33] René Ranftl, Alexey Bochkovskiy, and Vladlen Koltun. Vision transformers for dense prediction. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, pages 12179–12188, 2021.
- [34] René Ranftl, Katrin Lasinger, David Hafner, Konrad Schindler, and Vladlen Koltun. Towards robust monocular depth estimation: Mixing datasets for zero-shot cross-dataset transfer. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 44(3):1623–1637, 2022.
- [35] Onur Bagoren, Marc Micatka, Katherine A Skinner, and Aaron Marburg. Pugs: Perceptual uncertainty for grasp selection in underwater environments. In *2025 IEEE International Conference on Robotics and Automation (ICRA)*, pages 11278–11285. IEEE, 2025.
- [36] Nikhila Ravi, Valentin Gabeur, Yuan-Ting Hu, Ronghang Hu, Chaitanya Ryali, Tengyu Ma, Haitham Khedr, Roman Rädle, Chloe Rolland, Laura Gustafson, et al. Sam 2: Segment anything in images and videos. *arXiv preprint arXiv:2408.00714*, 2024.
- [37] Lihe Yang, Bingyi Kang, Zilong Huang, Zhen Zhao, Xiaogang Xu, Jiashi Feng, and Hengshuang Zhao. Depth anything v2. *Advances in Neural Information Processing Systems*, 37:21875–21911, 2024.
- [38] Nikita Karaev, Iurii Makarov, Jianyuan Wang, Natalia Neverova, Andrea Vedaldi, and Christian Rupprecht. Cotracker3: Simpler and better point tracking by pseudo-labelling real videos. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 6013–6022, 2025.
- [39] Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. Learning transferable visual models from natural language supervision. In *International conference on machine learning*, pages 8748–8763. PMLR, 2021.