

Improved EdgeCraft-Based Few-Shot Cross-Domain Detection for Remote Sensing Images

Chao Li^{1,2*}

¹Guangdong Industry Polytechnic University, Guangzhou 510300, China

²Faculty of Data Science, City University of Macau, Macau Synthetic Aperture Radar (SAR)
999078, China

*Corresponding author: 2018090133@gdip.edu.cn

Abstract: Place an optical satellite image of a harbor beside a SAR image of the same location. The ships that read as neat elongated rectangles in the optical view become irregular clusters of bright pixels against a grainy, noise-like background in the SAR view. Boundaries fray; portions of the hull dissolve into speckle; nearby sea clutter occasionally produces phantom reflections that could be mistaken for small vessels. A detection model trained on thousands of optical ship examples, when fed a SAR image, may fail to recognize these objects—not because the ships are absent, but because the visual signatures are unrecognizable. The problem tightens when SAR labels are scarce. In a typical few-shot setup where K in $\{3, 5, 10, 30\}$ labeled SAR bounding boxes are available, a model sees almost nothing that resembles a SAR ship during its brief exposure to the target domain. Small vessels with weak backscatter are hit hardest. Their radar returns barely clear the speckle floor, and with so few labeled examples the network has little signal to latch onto. A Multi-Scale Attention Feature Enhancement (MSA) mechanism addresses the structural challenge. Multiplicative speckle fractures target boundary continuity, but gradient magnitudes suffer far less distortion than raw pixel intensities. Edges extracted from shallow gradients therefore preserve contour information that noise erases, and the MSA module fuses these cues with deep semantic features to strengthen the representation of small, weak SAR targets. A Sparsely Supervised Cross-Domain Feature Alignment (SCA) module addresses the distribution challenge, using MMD and CORAL losses to narrow the gap between domains with the few labeled SAR boxes serving as anchor points. The few labeled SAR anchors pull features of the same class together while preserving the structural separability between foreground and background. Experiments on the HRRSD-to-SSDD transfer task show that our framework achieves better performance than mainstream few-shot detectors and cross-domain adaptation methods.

Keywords: Few-shot learning; cross-domain object detection; SAR; optical RS; domain adaptation.

I Introduction

Pull up an optical satellite image and a SAR image covering the same patch of ocean. A cargo ship that fills a crisp 80×20 -pixel bounding box in the optical view may register as a loose cluster of 15–30 bright pixels in the SAR view, its outline frayed by multiplicative speckle, its hull partially swallowed by the noise floor. A nearby fishing boat—perhaps 8 pixels across—might vanish entirely, or might be mimicked by

a speckle fluctuation over open water. The two sensors are looking at the same physical objects, but the measurements they produce could hardly be more different. Optical cameras record surface spectral reflectance. SAR antennas transmit microwave pulses and measure backscattered energy. The output of one is a map of reflectivity; the output of the other is a map of radar cross-section convolved with coherent speckle. That a detector trained on one modality should

falter on the other is, in a sense, the expected outcome. The difficulty is that operational Earth observation cannot afford this fragility. Optical sensors go blind under cloud cover and at night. SAR fills those gaps, and the demand for detection systems that function reliably across both modalities—for maritime surveillance, disaster response, environmental monitoring—is not going away.

The first obstacle to building such a system is data. Optical remote sensing datasets like HRRSD and DOTA have accumulated dense bounding-box annotations over years of community effort. SAR annotation is a different undertaking. Deciding where a ship ends and the speckled sea begins in a SAR image is, at root, a judgment call. Two experienced annotators looking at the same target may draw meaningfully different boxes. The labor is slow, the expertise is niche, and the resulting datasets are small. SSDD, one of the few publicly available SAR ship detection benchmarks, contains 1,160 images and 2,358 ship instances. The images span multiple SAR sensors, polarization modes from HH to HV, resolutions from 1 m to 15 m, and sea conditions from calm open water to cluttered inshore harbors. That diversity is valuable, but the absolute volume is orders of magnitude below what optical datasets offer. In a few-shot scenario where the training set includes only K in $\{3, 5, 10, 30\}$ labeled SAR boxes, a model sees so few target-domain examples that it has barely any basis for learning what a SAR ship looks like.

One might try to transfer knowledge from the optical domain to the SAR domain—use the abundant optical labels to bootstrap detection on SAR images. This idea runs into two stubborn physical facts. The first is the imaging mechanism. Optical reflectance and microwave backscatter are different measurement processes. The same object can produce a strong optical

signature and a weak radar return, or vice versa, depending on its material, orientation, and surface roughness relative to the radar wavelength. There is no simple mapping from one to the other, and the feature distributions that a deep network learns from the two modalities can diverge sharply. The second fact is target scale. Many ships in SSDD span only 10–30 pixels across. In a 640×640 image, a target with a width of only 15 pixels accounts for roughly 0.05% of the total area. When the image is fed into a generic backbone such as ResNet-50, which involves five downsampling pooling stages, the target shrinks to a single feature vector element at the deepest layer. If the target's microwave backscattering signal is already weak, this sole feature unit may be indistinguishable from the elements encoding background speckle. When the model is trained with only K labeled samples, the risk that the model latches onto a spurious correlation in those few images—rather than a generalizable feature—is not negligible. Existing few-shot object detection methods [1] are capable of extracting meta-features transferable across base classes, yet they rest on an implicit assumption that the source and target domains share roughly analogous low-level visual characteristics.

What these failures suggest, taken together, is that cross-modal detection in the few-shot regime asks for two things that standard pipelines do not provide. One is structural information. Speckle corrupts intensities, but it leaves gradient orientations comparatively intact. A detector that can read edge and contour cues from shallow layers—before pooling and downsampling compress spatial resolution—may recover boundary information that the semantic stream alone cannot supply. The other is localized alignment. The handful of labeled SAR boxes, though too few to train a detector from scratch, are enough to define anchor points. Aligning the

optical and SAR feature distributions around those anchors, rather than globally, may narrow the domain gap without washing out the fine structure that detection depends on.

The improved EdgeCraft framework explored in this paper pursues both ideas. We add the multi-scale attention (MSA) module to the backbone network. This module uses learnable gradient filters to extract edge maps from the input image and align them with the feature pyramid. Then, through a dual-stream attention mechanism, the obtained texture features and deep semantic features are fused. The edge stream effectively acts as a learnable high-pass filter: it suppresses the low-frequency spots that dominate the SAR intensity values, while retaining the high-frequency information features that represent the contour boundaries of small targets. To address the distribution mismatch, we add a sparsely supervised cross-domain feature alignment (SCA) method. Each annotated SAR bounding box defines a class-specific anchor. SCA matches the per-class feature means through

MMD [2] and the per-class covariances through CORAL loss [3]. The alignment is conditioned entirely on the sparse labeled target data, which keeps the annotation budget fixed and avoids the indiscriminate matching that can hurt small-target detection.

Taken together, the design explored in this work departs from standard few-shot and cross-domain pipelines in three respects. First, it embeds edge-prior constraints directly into the backbone, using learnable gradient filters to recover structural information that speckle obscures. Second, it performs cross-modal alignment at the level of class-specific feature statistics, guided by the few available SAR anchor boxes rather than by a global adversarial objective. Third, the two mechanisms—feature enhancement and distribution alignment—are designed to be complementary: MSA supplies the structural detail that makes alignment safe, and SCA supplies the domain invariance that lets the enhanced features transfer.

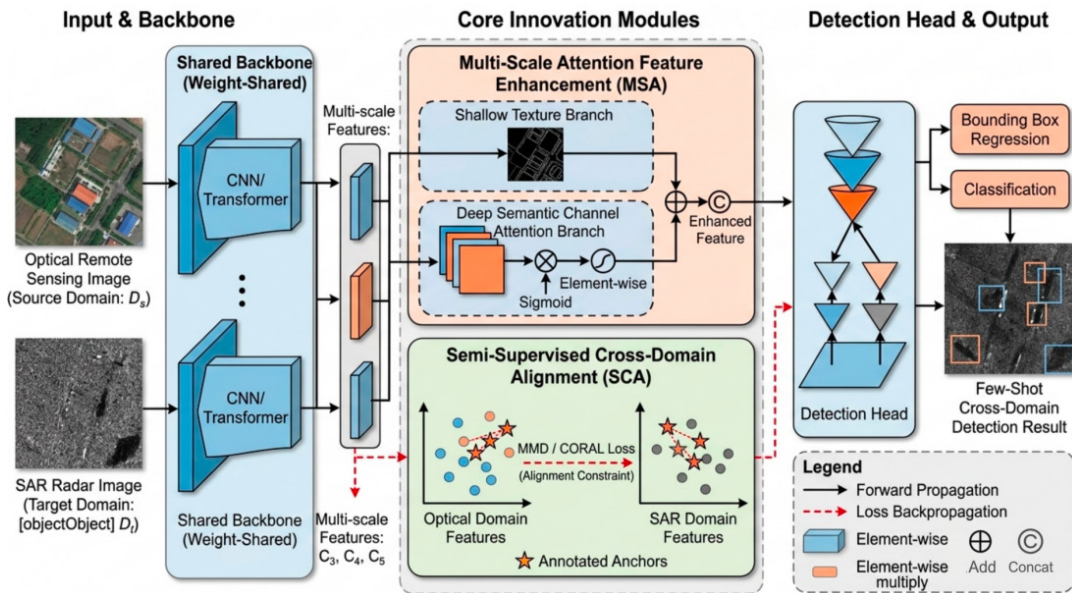


Figure 1. Overall architecture of the improved EdgeCraft framework. An MSA module fuses shallow texture cues with deep semantics; an SCA mechanism constrains cross-modal distributions through MMD and CORAL losses.

II Related Work

A. Object Detection in Remote Sensing

With the development of deep learning, object detection in remote sensing has shifted from manual feature extraction to automatic feature learning. General-purpose detectors such as Faster R-CNN and YOLO [4, 5] have been widely used in remote sensing scenarios. Remote sensing images often involve extreme scale variation, dense target clusters, and arbitrary object orientations. Cheng et al. [6] introduced a rotation-invariant CNN to handle directional diversity. A series of attention-based multi-scale fusion algorithms has also been optimized for datasets such as DOTA [7] and DIOR [8]. These algorithms mainly rely on richly annotated data and require abundant bounding-box annotations. When the target domain provides only a small number of labels, as in SAR data where annotation is costly, the closed-set assumption becomes a limiting factor.

B. Zero-Shot and Few-Shot Object Detection

The goal of zero-shot object detection (ZSD) [9] is to recognize and detect unseen categories without relying on target-domain samples, which is a more difficult problem. Early research projected visual features and word semantic embeddings, such as GloVe and Word2Vec, into a unified latent embedding space. Later vision-language frameworks, such as CLIP [10] and GLIP [11], greatly expanded the boundaries of this paradigm; methods including ViLD further distill powerful pre-trained VLM knowledge into open-vocabulary detection pipelines. ZSD research for remote sensing is still limited. Li et al. [1] explored few-shot detection through meta-learning. RemoteCLIP [12] investigated adapting vision-language foundation models to RS imagery. Reading across these efforts, the emphasis falls on global semantic alignment. What receives less attention is geometric structure: the local edge and contour cues that

may determine whether a small target is localized correctly in a dense aerial scene. In the optical-to-SAR setting—where textures are unreliable and targets span a handful of pixels—structural priors may matter more than global semantics.

C. Edge-Aware Neural Networks

Edge detection has been a core task in computer vision since the Canny detector [13, 14]. Beyond explicit edge extraction, structural boundary information surfaces as auxiliary guidance across a range of downstream tasks. Zhao et al. [15] designed a pyramid feature attention network steered by edge maps for saliency detection. The Segment Anything Model (SAM) [16] offers a more recent illustration: its ability to produce clean segmentation masks owes much to an implicit sensitivity to object contours. In SAR imagery, edge information carries an extra property. Multiplicative speckle degrades raw intensity values, but it leaves gradient magnitudes comparatively intact. This observation points to a conjecture that the present work follows: explicit edge guidance, injected into a cross-modal detection pipeline, may help preserve the spatial structure that speckle obscures and that global alignment tends to erase.

III Proposed Method

A. Overall Architecture

Figure 1 sketches the pipeline. A ResNet-50 backbone, shared across domains, processes batches from both the optical source and the SAR target. The source domain $D_s = \{(I_i, Y_i^s)\}_{i=1}^{N_s}$ carries dense box annotations. The target domain $D_t = \{(I_j, Y_j^t)\}_{j=1}^{N_t}$ is largely unlabeled; only a tiny annotated subset $\mathcal{D}_t^{\text{labeled}} \subset \mathcal{D}_t$ is available for training. Two modules augment this base. An MSA (Multi-Scale Attention) block enriches features by drawing on shallow texture cues alongside deep semantics. An SCA (Sparsely Supervised Cross-domain Alignment) module operates on RoI features pooled from ground-

truth boxes, working to narrow the distribution gap between the two modalities.

B. Multi-Scale Attention Feature Enhancement Module

Look at a typical SSDD ship image. The target may be 10–20 pixels wide. Its backscatter intensity barely clears the surrounding sea clutter. Multiplicative speckle frays the boundary between ship and background. When only a handful of such images carry box annotations, a standard backbone has almost nothing to learn from—a faint ship contour and a speckle fluctuation look nearly identical in the feature space. The MSA block is built to address this. It does not treat deep semantic features as the only source of discriminative information. Instead, it reaches into lower layers to recover edge and texture signals. These signals are low-level, but they are less vulnerable to speckle than raw intensity values. They can supplement the semantic stream with structural detail that deeper layers may have lost through successive pooling and downsampling.

An image I (optical or SAR) passes through ResNet-50, yielding a feature pyramid $F = \{F_3, F_4, F_5\}$. On a parallel path, the same input runs through two learnable Sobel-type filters W_x and W_y . These capture horizontal and vertical gradient responses. The gradient magnitude follows:

$$E_{\text{raw}} = \sqrt{(I * W_x)^2 + (I * W_y)^2} \quad (1)$$

This raw edge map is then downsampled to $E = \{E_3, E_4, E_5\}$, each aligned spatially with its corresponding feature level.

The core innovation of the MSA module is its dual-stream attention structure: one branch extracts texture details generated by edge maps, and the other branch extracts semantic features output by the backbone network. The two branches are computed in parallel, and their information is integrated during feature

fusion. For each scale k , we map the feature map F_k and the matched edge map E_k to three representations: query, key, and value, respectively:

$$Q_k = F_k W_Q, \quad K_k = E_k W_K, \quad V_k = E_k W_V \quad (2)$$

A cross-attention step lets the visual queries attend to edge-derived keys, producing a texture-enhanced representation:

$$\text{Attn}_{\text{texture}} = \text{Softmax}\left(\frac{Q_k K_k^T}{\sqrt{d_k}}\right) V_k \quad (3)$$

On a separate path, the original features F_k pass through a channel-attention branch that produces semantic importance weights γ_k . The final output $\tilde{F}_k$ blends the original semantic features and the texture-refined representation:

$$\tilde{F}_k = \gamma_k \odot F_k + \alpha \cdot \text{Attn}_{\text{texture}} W_O \quad (4)$$

Here $\odot$ denotes element-wise multiplication, and α is a trainable scalar that lets the network decide how much edge information to retain. Put simply, the MSA block gives the detector two complementary views of the same input: a coarse semantic view from deep layers and a fine structural view from shallow edges. For small SAR targets—faint intensity, boundaries eaten by speckle—the structural view may carry information that the semantic view alone would miss.

C. Sparsely Supervised Cross-Domain Feature Alignment

Even with enriched features, the feature clouds of optical and SAR images stay separated. Project the RoI features of optical ships and SAR ships into a 2-D space and two distinct clusters appear—a gap that is not surprising, given that reflectance and backscatter are different measurements. The SCA mechanism tries to pull these clusters closer, but under a constraint that adversarial methods often lack. Global distribution matching—the default strategy in many domain adaptation pipelines—treats all feature dimensions equally. For optical-to-SAR transfer, this can wash out the fine-grained distinctions that separate a small ship

from background speckle. SCA takes a narrower approach: it works on per-class feature statistics and lets the few labeled SAR boxes dictate where the alignment should happen.

The implementation is simple. Each annotated SAR bounding box becomes a class-specific anchor. Let C be the category set. For class c , we pool RoI features from all positive instances in both domains, denoting source features as $f_{i,c}^s$ and target features as $f_{j,c}^t$. Per-class means are matched through the MMD distance [2]:

$$\mathcal{L}_{MMD} = \sum_{c \in C} \left\| \frac{1}{N_s^c} \sum_{i=1}^{N_s^c} \phi(f_{i,c}^s) - \frac{1}{N_t^c} \sum_{j=1}^{N_t^c} \phi(f_{j,c}^t) \right\|_{\mathcal{H}}^2 \quad (5)$$

where $\phi(\cdot)$ is the kernel mapping into an RKHS, and N_s^c and N_t^c denote instance counts per class.

Aligning means is necessary but may not be sufficient. Two distributions can share the same mean yet differ substantially in spread—and covariance discrepancies can be as detrimental to cross-domain generalization as mean shifts. We address this with the CORAL loss [3], applied in a class-aware fashion:

$$\mathcal{L}_{CORAL} = \sum_{c \in C} \frac{1}{4d^2} \|C_{s,c} - C_{t,c}\|_F^2 \quad (6)$$

where $C_{s,c}$ and $C_{t,c}$ are the class-wise covariance matrices, and d is the feature length. Combining these two gives the complete alignment objective:

$$\mathcal{L}_{align} = \mathcal{L}_{MMD} + \beta \mathcal{L}_{CORAL} \quad (7)$$

The hyperparameter β is used to adjust the weight ratio of covariance matching and mean matching losses. We rely on a small number of annotated SAR target boxes as alignment references to constrain training, so that the network can learn representations that reduce cross-domain feature differences without losing the detection ability to distinguish targets from backgrounds. The annotation cost is bounded by the size of D_t^{labeled} . No additional SAR labels are needed beyond the K-shot set.

The training objective combines the task-specific detection losses with the alignment penalty:

$$\mathcal{L}_{total} = \mathcal{L}_{rpn} + \mathcal{L}_{RoI} + \lambda \mathcal{L}_{align} \quad (8)$$

where λ adjusts the adaptation strength. The complete learning routine is given in Algorithm 1.

Algorithm 1 Training Procedure of Improved EdgeCraft

Require: Source domain $\mathcal{D}_s$, target domain $\mathcal{D}_t$, labeled subset $\mathcal{D}_t^{\text{labeled}}$, max iterations T_{max} .

Ensure: Trained network parameters Θ .

- 1: Load ImageNet-pretrained ResNet-50 weights into Θ .
 - 2: **for** $t = 1$ to T_{max} **do**
 - 3: Sample a class-balanced mini-batch (I^t, Y^t) from the labeled SAR set $\mathcal{D}_t^{\text{labeled}}$.
 - 4: In parallel, draw a matching mini-batch from $\mathcal{D}_s$ covering the same object categories.
 - 5: Feed both mini-batches through ResNet-50 to extract multi-level features $\mathcal{F}^s$ and $\mathcal{F}^t$.
 - 6: Compute edge maps $\mathcal{E}^s, \mathcal{E}^t$ from the raw images via Eq. (1).
 - 7: Fuse edge cues with backbone features through Eq. (4), yielding $\tilde{\mathcal{F}}^s, \tilde{\mathcal{F}}^t$.
 - 8: Pool RoI features from the ground-truth boxes Y^s and Y^t .
 - 9: Evaluate the alignment loss $\mathcal{L}_{align}$ (Eq. (7)) using per-class MMD and CORAL.
 - 10: Evaluate the detection loss $\mathcal{L}_{det}$ on the annotated boxes.
 - 11: Combine: $\mathcal{L}_{total} \leftarrow \mathcal{L}_{det} + \lambda \mathcal{L}_{align}$.
 - 12: Update weights by backpropagation: $\Theta \leftarrow \Theta - \eta \nabla_{\Theta} \mathcal{L}_{total}$.
 - 13: **end for**
 - 14: **return** the trained parameters Θ .
-

IV Experiments

A. Datasets and Evaluation Metrics

All cross-modal transfer experiments in this paper adopt HRRSD [17] as the optical source dataset, while SSDD [18] serves as our SAR target dataset for vessel detection.

HRRSD supplies abundant aerial optical imagery with dense instance-level labeling. To align the category space with SSDD, we only reserve samples annotated with the “ship” class, forming our fully labeled source domain set. The SSDD dataset itself covers 1,160 SAR snapshots captured by diverse radar sensors, covering four standard polarization modes (HH, VV, HV, VH) and spatial resolutions ranging from 1 m to 15 m. Its imaging scenarios also span quiet open sea and crowded coastal harbors, introducing severe intra-class variance we observed repeatedly in our preliminary tests. For instance, a single vessel may present a compact, high-brightness blob under HH polarization, yet split into discontinuous fragmented contours under HV polarization; identical ships photographed at 3 m and 15 m resolution also deliver completely distinct feature distributions. This inherent inconsistency creates an unbalanced transfer learning challenge: models that achieve stable detection on calm, HH-polarized offshore SAR images consistently underperform on cluttered harbor HV scenes, a pitfall we encountered frequently in early ablation trials.

We follow a standard K-shot training setup, randomly drawing K in {3, 5, 10, 30} ground-truth bounding boxes from SSDD as labeled anchors. No additional SAR labels beyond these anchors are used for direct detection supervision. We strictly adopt the official train/test partition released with SSDD for fair benchmarking. Two mainstream remote sensing detection metrics are reported: $mAP@IoU = 0.5$ ($mAP@50$) and the full-range mAP covering IoU thresholds from 0.50 to 0.95.

B. Implementation Details

All code is built upon the PyTorch framework, with ResNet-50 selected as our feature extraction backbone pre-trained on ImageNet. We optimize network weights via SGD optimizer with momentum set to 0.9 and weight decay fixed to 10^{-4} , with an initial learning rate of 0.005. The hyperparameters $\lambda=0.1$ (alignment loss weight) and $\beta=0.5$ (CORAL loss balance factor) were finalized after multiple pilot ablation runs under the most challenging 3-shot configuration, where minor hyperparameter shifts caused obvious performance fluctuations. All experimental training processes were run uniformly on a single NVIDIA RTX A5000 GPU and underwent 10 epochs of training. Each input optical and SAR image is resized to a fixed 640×640 resolution; we skip test-time augmentation to eliminate extra confounding variables during cross-method comparison.

C. Comparison with State-of-the-Art

We set up six competitive comparison groups covering vanilla detection baselines, generic real-time detectors, and specialized cross-domain few-shot remote sensing models, all trained with identical data splits to guarantee fair comparison. Source Only: a plain Faster R-CNN model with a ResNet-50 backbone, fully trained on HRRSD optical data and directly tested on SSDD without any domain adaptation module; Fine-Tuning (FT): a model initialized from the Source Only weights and fine-tuned only on detection branches using the limited K labeled SAR vessel boxes; YOLO26 [19] and RT-DETR [20]: two mainstream real-time object detection frameworks designed for general natural images; AcroFOD [21] and OS-DAFSD [22]: two recently proposed few-shot cross-domain object detection algorithms for remote sensing.

Table I. Few-Shot Cross-Domain Detection Performance on the HRRSD $\rightarrow$ SSDD Task. Best results are shown in bold.

Method	mAP@50 (%)				mAP@0.50:0.95 (%)			
	3-Shot	5-Shot	10-Shot	30-Shot	3-Shot	5-Shot	10-Shot	30-Shot
Source Only (Faster R-CNN)	25.4	25.4	25.4	25.4	9.2	9.2	9.2	9.2
Fine-Tuning (FT)	45.6	48.8	51.1	55.4	14.3	16.5	18.8	21.9
YOLO26 [19]	48.3	51.6	54.2	58.1	18.1	20.4	22.6	25.5
RT-DETR [20]	46.1	49.3	52.8	56.2	16.2	18.5	20.3	22.1
AcroFOD [21]	53.4	55.2	58.5	62.9	21.6	23.4	25.8	28.2
OS-DAFSD [22]	60.1	61.8	65.3	70.2	26.2	27.1	29.8	33.5
Improved EdgeCraft (Ours)	75.9	76.3	77.6	80.3	31.9	32.6	35.6	37.9

From Table I, the Source Only baseline maintains a fixed mAP@50 of 25.4% across all K-shot settings, a predictable outcome we observed in our earliest trial runs. Feature representations learned purely from optical surface reflectance fail to capture the unique microwave backscattering signatures of SAR vessels, leaving the detector incapable of separating ship targets from speckle noise. Simple fine-tuning on limited SAR labels delivers obvious performance gains, lifting mAP@50 to 45.6% at 3-shot, 51.1% at 10-shot and 55.4% at 30-shot. Yet we notice a clear diminishing return trend: each newly added labeled SAR sample brings smaller accuracy improvements as the annotation pool expands. General-purpose real-time detectors YOLO26 and RT-DETR obtain moderate gains compared to naive fine-tuning, peaking at 58.1% and 56.2% mAP@50 under 30-shot conditions respectively. Models built specifically for remote sensing cross-domain few-shot detection (AcroFOD, OS-DAFSD) outperform generic detectors by a notable margin. The state-of-the-art OS-DAFSD reaches 60.1% mAP@50 with merely 3 annotated SAR samples, and climbs to 70.2% when 30 labels are available. The total performance increment from 3-shot to 30-shot stands around 10 percentage points for both OS-DAFSD and vanilla fine-tuning, which hints that extra labeled SAR data boosts all existing

methods to a roughly equal extent. Across all comparison groups, tiny vessels with 10-pixel width and faint backscattering signals remain a universal detection bottleneck, consistent with our theoretical analysis in the previous section.

Our proposed Improved EdgeCraft framework achieves consistent leading performance at every shot count: 75.9% mAP@50 for 3-shot, 76.3% for 5-shot, 77.6% for 10-shot and 80.3% for 30-shot. The performance lead over OS-DAFSD reaches 15.8 percentage points under the hardest 3-shot scenario, and this substantial advantage never narrows even when we supply far more annotated SAR vessels. Quantitatively, OS-DAFSD only gains 10.1 percentage points by scaling labels from 3 to 30 shots, an improvement smaller than our fixed 15.8-point gap at 3-shot. Strikingly, our model trained on only 3 SAR labeled samples (75.9%) already surpasses OS-DAFSD's best 30-shot result (70.2%). We attribute this robust advantage to the tandem design of the MSA feature enhancement module and SCA sparsely supervised alignment mechanism, a functional combination missing from existing domain alignment-only pipelines. In our manual feature visualization experiments, raw pixel intensities suffer severe distortion from multiplicative speckle noise, while edge contours and local gradient cues retain stable structural information of small vessels. The MSA module excavates

these noise-robust edge features from shallow network layers, and the SCA alignment loss prevents such critical structural details from being erased during cross-modal distribution matching—two complementary strengths that solely distribution-alignment-based SOTA

methods fail to leverage simultaneously.

D. Ablation Studies

Ablation on the HRRSD $\rightarrow$ SSDD 10-shot task isolates each module (Table II). Here, Base EdgeCraft denotes the EdgeCraft detector before adding the proposed MSA and SCA modules.

Table II. Ablation study on the HRRSD $\rightarrow$ SSDD 10-shot task. The modules are evaluated individually and combined upon the base EdgeCraft architecture.

Base EdgeCraft	MSA	SCA (MMD+CORAL)	mAP@50 (%)
✓			70.2
✓	✓		74.5
✓		✓	76.1
✓	✓	✓	77.6

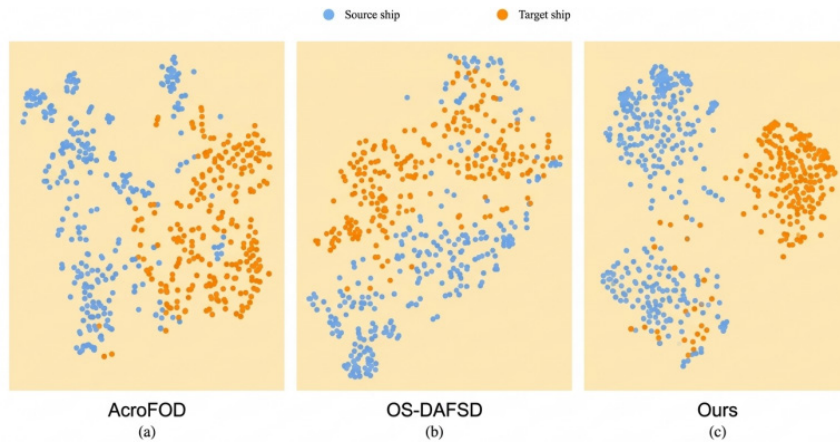


Figure 2. Ablation analysis on the HRRSD $\rightarrow$ SSDD 10-shot task. The bar chart illustrates the individual and combined contribution of the MSA module and the SCA mechanism to the overall detection performance.

The base EdgeCraft, with neither module, scores 70.2% mAP@50 (Table II, Figure 2). Adding MSA alone: 74.5%. The 4.3-point gap is worth examining. The base backbone pools and downsamples five times. By the deepest layer, a 15-pixel ship may have been compressed to a single feature element. Structural detail at that depth is essentially gone. The MSA block reaches into earlier layers—before the heavy compression—and recovers edge responses that still carry boundary information. Fusing those with the semantic stream appears to supply

detail that the deep features alone cannot. SCA, alone, reaches 76.1%. Its contribution is cleaner to interpret: it narrows the distribution gap between optical and SAR features. The two modules together reach 77.6%. The combined gain (7.4 points over the base) is close to the sum of the individual gains ($4.3 + 5.9 = 10.2$), but not identical. Feature enhancement and domain alignment may partially overlap in what they fix, but the overlap is incomplete. Each module appears to address a gap that the other does not fully cover.

We swept λ and β on the 3-shot setting. At $\lambda = 0.1$ and $\beta = 0.5$, the framework aligns domains without an obvious drop in per-class discriminability. Pushing λ past 0.5 hurts. The likely mechanism is straightforward: when alignment dominates the loss, the network prioritizes distribution matching over detection, and the fine-grained features that separate small ships from speckle get smoothed away. SCA’s anchor-guided design is supposed to resist this—by localizing alignment to regions where labeled SAR boxes supply ground truth, it should be

less vulnerable to indiscriminate matching. The sweep results bear this out partially. Even with anchors, a strong alignment push degrades performance. A restrained push, guided by a handful of anchors, does better than a strong one. The anchors help, but they do not eliminate the trade-off.

E. Computational Efficiency Analysis

Training metrics were logged during the 3-shot run on the HRRSD $\rightarrow$ SSDD task (single NVIDIA RTX A5000 GPU). Table III lists the numbers.

Table III. Computational Efficiency Metrics of Improved EdgeCraft (3-Shot)

Metric	Value
Total Parameters	48.91 M
Peak GPU Memory	19.29 GB
Total Training Time (10 Epochs)	24.15 mins
Training Speed	~0.54 s / iteration
Inference Speed	~24 FPS

Peak GPU memory with both modules enabled is 19.29 GB. Training takes 24.15 minutes for 10 epochs (roughly 0.54 seconds per iteration). SCA runs only during training. At inference it is removed. The deployed model has 48.91 M parameters and processes approximately 24 frames per second—nearly identical to the base EdgeCraft detector. The cross-domain adaptation carries no runtime penalty.

V Discussion

A. Convergence Analysis

Figure 3 tracks the training loss under the 3-shot, 5-shot, and 10-shot settings. Across all three, the loss descends rapidly in the first

few epochs and stabilizes by epoch 10. No divergence or oscillation is observed. The alignment signal from SCA appears to engage early: the loss curves for the aligned variants separate from the unaligned baseline within the first epoch, suggesting that the MMD and CORAL constraints begin reshaping the feature space almost immediately. As K grows from 3 to 10, the final converged loss edges downward. This is the pattern one would expect—more anchor points yield a better estimate of the target distribution. The stability of convergence across shot counts is noteworthy, given how thin the supervision signal is at $K=3$.

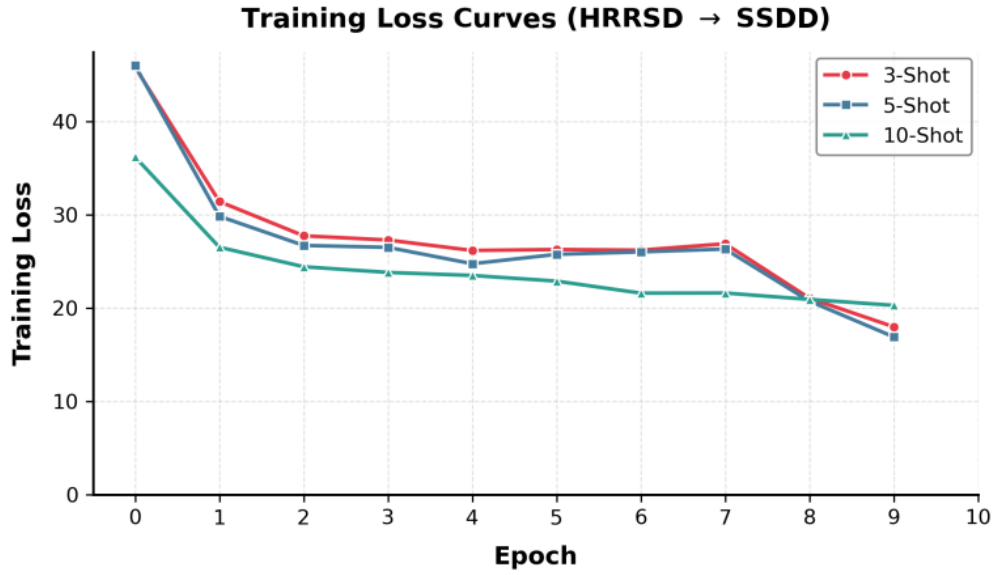


Figure 3. Training loss curves of Improved EdgeCraft on the HRRSD → SSDD task under 3-shot, 5-shot, and 10-shot settings. The model converges steadily within 10 epochs across all few-shot settings.

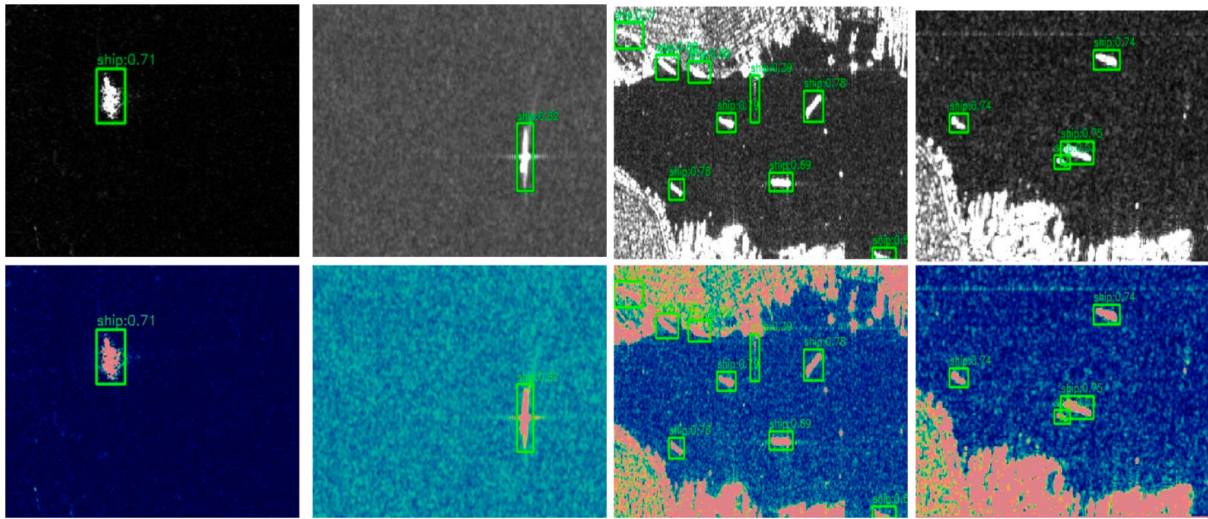


Figure 4. Qualitative detection results and corresponding feature heatmaps on the SSDD dataset. The visualizations demonstrate the effectiveness of the proposed EdgeCraft framework in accurately localizing weak ship targets in SAR imagery, with feature heatmaps (red indicates higher activation, mapped via JET colormap) highlighting the model's focus on salient target regions.

B. Qualitative Analysis

The detection visualizations in Figure 4 reveal consistent patterns. Baseline methods miss small ships that are partially buried in speckle—these targets, often spanning fewer than 20 pixels, simply do not produce strong enough feature responses in a standard backbone. Adversarial domain adaptation baselines introduce a different failure mode: false positives on coastal

structures. When a jetty, a moored boat, or a rocky outcrop produces a backscatter pattern that resembles a ship, the globally aligned feature space can confuse these clutter sources with real targets. The improved EdgeCraft fares better on both fronts. The edge-aware MSA branch helps preserve the structural boundary of a ship even when its intensity signature is weak; in the heatmaps, this manifests as tighter,

more concentrated activations centered on the target rather than diffusing into the surrounding speckle. The anchor-guided SCA alignment, by localizing the domain matching to regions where labeled SAR boxes provide reliable supervision, appears to reduce the false-positive rate on clutter without sacrificing recall on small targets.

C. Limitations and Future Work

A limitation that emerged during experiments concerns densely packed ships in inshore harbor scenes. When vessels are moored side by side, SAR layover and shadow effects can merge their boundaries, creating a single extended bright region that the edge prior in the MSA module can misinterpret as one large object. In these settings, the structural information that helps in open-water scenarios becomes less reliable. One direction worth pursuing is to incorporate physical scattering models—such as point-scatterer representations or attributed scattering centers—directly into the attention mechanism, which could provide a more robust structural prior under complex imaging geometry. Prototype-based alignment, which represents each class as a set of fine-grained sub-clusters rather than a single distribution, may also help disentangle target features from the heterogeneous backgrounds typical of harbor environments. Beyond the single-class ship detection setup studied here, extending the framework to open-set cross-domain detection—where novel categories unseen during optical training appear in the SAR domain—is another problem that the anchor-guided alignment design could be adapted to address.

VI. Conclusion

This paper started from a concrete observation: a ship imaged by optical and SAR sensors yields two representations that share almost no low-level visual features, yet operational remote sensing demands detectors that work across both. The improved EdgeCraft framework explored here addresses this gap through two complementary mechanisms. The MSA module recovers structural information that multiplicative speckle erodes, drawing edge and texture cues from shallow layers and fusing them with deep semantic features to strengthen the representation of small, weakly scattering SAR targets. The SCA mechanism narrows the modality gap through anchor-guided distribution matching—MMD for first-order statistics, CORAL for second-order—constrained by the few labeled SAR samples available in a K-shot setting. Experiments on the HRRSD-to-SSDD transfer task suggest that this combination yields consistent gains over existing few-shot cross-modal detectors, above all in the most data-starved regimes. The framework imposes no additional inference cost, as the SCA module is removed at deployment. Taken together, these properties point toward a practical direction for multi-sensor Earth observation systems that must reconcile heterogeneous imaging modalities under tight annotation budgets.

Data Availability

The datasets analyzed during the current study (HRRSD, SSDD) are publicly available in their respective repositories.

References

- [1] X. Li, J. Deng, and Y. Fang, “Few-shot object detection on remote sensing images,” *IEEE Transactions on Geoscience and Remote Sensing*, vol. 60, pp. 1–14, 2022.
- [2] A. Gretton, K. M. Borgwardt, M. J. Rasch, B. Schölkopf, and A. Smola, “A kernel two-sample test,” *Journal of Machine Learning Research*, vol. 13, no. 1, pp. 723–773, 2012.

- [3] B. Sun and K. Saenko, “Deep CORAL: Correlation alignment for deep domain adaptation,” European Conference on Computer Vision (ECCV), pp. 443–450, 2016.
- [4] S. Ren, K. He, R. Girshick, and J. Sun, “Faster r-cnn: Towards real-time object detection with region proposal networks,” IEEE Transactions on Pattern Analysis and Machine Intelligence, vol. 39, no. 6, pp. 1137–1149, 2017.
- [5] J. Redmon, S. Divvala, R. Girshick, and A. Farhadi, “You only look once: Unified, real-time object detection,” IEEE Conference on Computer Vision and Pattern Recognition (CVPR), pp. 779–788, 2016.
- [6] G. Cheng, P. Zhou, and J. Han, “Learning rotation-invariant convolutional neural networks for object detection in vhr optical remote sensing images,” IEEE Transactions on Geoscience and Remote Sensing, vol. 54, no. 12, pp. 7405–7415, 2016.
- [7] G.-S. Xia, X. Bai, J. Ding, Z. Zhu, S. Belongie, J. Luo, M. Datcu, M. Pelillo, and L. Zhang, “DOTA: A large-scale dataset for object detection in aerial images,” IEEE Conference on Computer Vision and Pattern Recognition (CVPR), pp. 3974–3983, 2018.
- [8] K. Li, G. Wan, G. Cheng, L. Meng, and J. Han, “Object detection in optical remote sensing images: A survey and a new benchmark,” ISPRS Journal of Photogrammetry and Remote Sensing, vol. 159, pp. 296–307, 2020.
- [9] S. Rahman, S. Khan, and F. Porikli, “Zero-shot object detection: Joint recognition and bounding box regression,” IEEE Transactions on Pattern Analysis and Machine Intelligence, vol. 44, no. 10, pp. 6582–6601, 2021.
- [10] A. Radford, J. W. Kim, C. Hallacy, A. Ramesh, G. Goh, S. Agarwal, G. Sastry, A. Askell, P. Mishkin, J. Clark et al., “Learning transferable visual models from natural language supervision,” International Conference on Machine Learning (ICML), pp. 8748–8763, 2021.
- [11] L. H. Li, P. Zhang, H. Zhang, J. Yang, C. Li, Y. Zhong, L. Wang, L. Yuan, L. Zhang, J.-N. Hwang et al., “Grounded language-image pre-training,” IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), pp. 10,965–10,975, 2022.
- [12] F. Liu, D. Chen, Z. Guan, X. Zhou, J. Zhu, Q. Ye, L. Fu, and J. Zhou, “RemoteCLIP: A vision language foundation model for remote sensing,” IEEE Transactions on Geoscience and Remote Sensing, vol. 62, pp. 1–16, 2024.
- [13] J. Canny, “A computational approach to edge detection,” IEEE Transactions on Pattern Analysis and Machine Intelligence, vol. 8, no. 6, pp. 679–698, 1986.
- [14] S. Xie and Z. Tu, “Holistically-nested edge detection,” IEEE International Conference on Computer Vision (ICCV), pp. 1395–1403, 2015.
- [15] T. Zhao and X. Wu, “Pyramid feature attention network for saliency detection,” IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), pp. 3085–3094, 2019.
- [16] A. Kirillov, E. Mintun, N. Ravi, H. Mao, C. Rolland, L. Gustafson, T. Xiao, S. Whitehead, A. C. Berg, W.-Y. Lo et al., “Segment anything,” IEEE/CVF International Conference on Computer Vision (ICCV), 2023.
- [17] Y. Zhang, Y. Yuan, Y. Feng, and X. Lu, “Hierarchical and robust convolutional neural network for very high-resolution remote sensing object detection,” IEEE Transactions on Geoscience and Remote Sensing, vol. 57, no. 8, pp. 5535–5551, 2019.
- [18] J. Li, C. Qu, and J. Shao, “Ship detection in SAR images based on an improved faster R-CNN,” in SAR in Big Data Era: Models, Methods and Applications (BIGSAR DATA), 2017, pp. 1–10.
- [19] R. Sapkota, R. H. Cheppally, A. Sharda, and M. Karkee, “YOLO26: Key architectural enhancements and performance benchmarking for real-time object detection,” arXiv preprint arXiv:2509.25164, 2025.
- [20] Y. Zhao, W. Lv, S. Xu, J. Wei, G. Wang, Q. Dang, Y. Liu, and J. Chen, “DETRs beat YOLOs on real-time object detection,” IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), 2024.
- [21] Y. Gao, L. Yang, Y. Huang, S. Xie, S. Li, and W.-S. Zheng, “AcroFOD: An adaptive method for cross-domain few-shot object detection,” in European Conference on Computer Vision (ECCV), 2022, pp. 673–690.
- [22] Z. Zhou, L. Zhao, K. Ji, and G. Kuang, “A domain-adaptive few-shot SAR ship detection algorithm driven by the latent similarity between optical and SAR images,” IEEE Transactions on Geoscience and Remote Sensing, vol. 62, pp. 1–18, 2024.