



iJRASET

International Journal For Research in
Applied Science and Engineering Technology



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Volume: 14 **Issue:** VII **Month of publication:** July 2026

DOI: <https://doi.org/10.22214/ijraset.2026.84392>

www.ijraset.com

Call: ☎ 08813907089

E-mail ID: ijraset@gmail.com

A Comparative Analysis of AI-Driven Marine Debris Detection Methods for Indian Ocean Plastic Monitoring

Nimisha Nair¹, Darshan Gadekar², Nalaksh Randhawa³, Roshan Kotkondawar⁴, Krushna Taiwade⁵
Department of Computer Engineering, St. Vincent Pallotti College of Engineering and Technology, Nagpur, India

Abstract: Plastic pollution continues to accumulate in marine environments at an alarming rate, with millions of tonnes entering the oceans annually. Given the length of India's coastline, early detection of floating debris is essential, as undetected objects may disperse, fragment, or sink before any cleanup effort can be initiated. Satellite remote sensing paired with artificial intelligence has become the go-to tool for this problem, yet a close reading of the published work reveals something surprising: researchers rarely explain why they chose one detection method over another. Most studies simply use whatever technique fits the sensor already in hand, without weighing it against alternatives suited to the deployment at hand — a gap even more pronounced for Indian coastal waters, which remain strikingly under-studied relative to the wider global literature. Motivated by this, the present paper reviews twenty-five published studies and compares seven AI-driven detection-method families — spectral-index screening, classical machine learning, CNN and U-Net segmentation, YOLO-based detection, hyperspectral classification, SAR-based screening, and drift-forecasting models — not merely on accuracy, but on sensor cost, robustness to confounders such as algae and biofouling, and real-world deployment readiness. The review finds that spectral-index and classical machine-learning approaches remain the most field-tested and affordable, reporting 86–98% accuracy on Sentinel-2 imagery, while drift forecasting is still reported largely qualitatively, with little validation against real drift tracks, and Indian coastal research remains scattered, single-site, and hard to compare across studies. These findings point to a pressing need for low-cost, transparently reported detection frameworks built for Indian shores, alongside a coordinated regional benchmark, drift-forecasting pipelines validated against in-situ tracks, and multi-sensor fusion combining optical, hyperspectral, and SAR data.

Keywords: marine debris detection; Floating Debris Index; Sentinel-2; MARIDA; YOLO; U-Net; hyperspectral imaging; Lagrangian drift modelling; Indian Ocean; machine learning.

I. INTRODUCTION

Millions of tonnes of plastic slip into the ocean every year, and along a coastline as vast as India's, the difference between a piece of debris that gets caught and one that doesn't often comes down to a matter of days: once missed, it simply keeps moving, drifting, fragmenting, or sinking beyond recovery [1]. Marine debris monitoring today rests on three pillars: field surveys conducted on foot or by boat, close-range platforms such as drones and unmanned surface vehicles, and satellite remote sensing. Among these, only satellites offer continuous, basin-scale visibility that no field team could replicate — but that promise is only realized once raw imagery is converted into something a system can act on: a spectral index, or a feature learned by a model [1][2].

That conversion step is where the field truly diverges. Turning a satellite pixel into a debris label depends entirely on the detection method chosen, and options abound spectral indices [1][2][11], classical machine learning [6][16]–[18], CNN and U-Net segmentation [3]–[5][13][24][25], YOLO-style detectors [8][12], hyperspectral classifiers [10][14], SAR-based screening [7][24], and an emerging class of drift-forecasting models that project where detected debris is headed next [9][15]. Each carries its own trade-offs in resolution, sensor cost, and deployment readiness. Yet across the literature, one gap stands out: these choices are rarely justified. Most studies default to whatever method fits the sensor already at hand, with little systematic weighing of alternatives against the deployment they are meant to serve. This paper confronts that gap head-on. It first synthesizes the published literature on these detection-method families, comparing them not on benchmark leaderboards alone, but on the criteria that matter for continuous, agency-facing coastal monitoring. It then uses that synthesis to *justify* — rather than merely describe — the design of an implemented decision-support pipeline built around Floating Debris Index screening [1], CNN-based classification [4][5][24], density-based hotspot clustering [13], and drift forecasting [9][15].

The case study is reported strictly at the architecture level: no externally validated field-accuracy figures are claimed, and every stage is labelled honestly as either live or running on simulated fallback data — a distinction Section V revisits in detail.

Four contributions follow: (i) a structured comparison of seven detection-method families, rather than a single-sensor summary; (ii) a sharper picture of where the literature falls short — chiefly around validated drift forecasting, the thin representation of Indian coastal studies, and vague reporting of what is actually live versus simulated in deployed systems; (iii) a literature-backed justification for the specific method combination used in the case study; and (iv) an architecture offered as an illustration of these gaps, not a claim to have closed them.

II. RELATED WORK

General reviews of remote-sensing-based marine debris detection catalog the field's progression from handcrafted spectral indices to deep learning and consistently flag sub-pixel targets and spectral confusion with algae or foam as unresolved challenges, yet few treat drift forecasting and detection as parts of a single pipeline rather than two separate literatures [3][21]. The most widely used open benchmark, MARIDA, reports that a random-forest baseline ($F1 \approx 0.70$) outperformed a U-Net baseline ($F1 \approx 0.50$) on the same weakly supervised segmentation task. This result is easy to over-read as "classical methods beat deep learning," but more plausibly reflects the benchmark's limited annotated volume relative to what deep segmentation networks typically need [3].

Within the detection literature, a handful of methods have been reported. Spectral-index screening based on the Floating Debris Index (FDI) underlies some of the earliest and most-cited satellite pipelines [1][2]. CNN and U-Net segmentation have become the default for satellite-scale mapping once labelled data are available [3]-[5][13][24][25], while YOLO-family detectors dominate close-range UAV and USV imagery where inference latency matters more than segmentation granularity [8][12]. Hyperspectral classifiers appear less often, largely because hyperspectral sensors remain more expensive and lower-throughput than multispectral ones; however, they report markedly finer material discrimination [10][14]. Notably, the justification for choosing one family over another in this literature is almost always a single sentence citing sensor availability rather than a structured comparison against alternatives — a pattern this paper tries to address directly, and a pattern that also shows up, distinctly, in the Indian coastal literature reviewed in Section III [16]-[20][22].

III. COMPARATIVE ANALYSIS OF MARINE DEBRIS DETECTION METHODS

Fig. 1 situates the reviewed method families along an approximate release timeline, from the first operational spectral index to today's hyperspectral and hybrid detection-plus-drift systems. Each family is discussed below and compared in Table I.

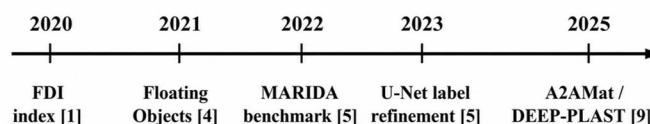


Figure 1 Approximate evolution of marine debris detection methods relevant to this study, from the first operational spectral index through hybrid detection-plus-drift systems.

A. Spectral-Index and Classical Machine-Learning Screening

Biermann et al. introduced the FDI and paired it with a naïve Bayes classifier, reporting a classification accuracy of approximately 86% on Sentinel-2 patches under controlled coastal conditions [1]. Themistocleous et al. developed a companion Plastics Index from field spectroradiometer measurements, finding the near-infrared band to be the most reliable discriminator [2]. Sannigrahi et al. layered random forests and support-vector machines on top of FDI, PI, and NDVI features, reporting balanced accuracies of roughly 80–98% (SVM) and 87–97% (RF) during model development, and 91% accuracy when the final model was tested on real-world data from Calabria and Beirut [6] — a range that overlaps with, rather than clearly exceeds, Biermann's 86%. Kremezi et al. tackled Sentinel-2's core limitation directly by fusing 10 m bands with 2 m World View imagery, pushing the smallest reliably detectable target down to approximately 0.6 m² [7]. This sub-strand is cheap and interpretable, but its accuracy is visibly threshold- and site-dependent, and none of these studies fully separate plastic from spectrally similar organic matter without a supervised stage layered on top.

B. CNN and U-Net Segmentation

Kikaki et al.'s MARIDA benchmark remains the standard reference point here, and its counterintuitive finding — random forest edging out a baseline U-Net — is discussed above [3]. Mifdal et al. released a complementary floating object dataset with pretrained CNN baselines aimed at global rather than site-specific coverage [4]. Rußwurm et al. then showed that correcting label noise in these same datasets, rather than deepening the network, produced most of the subsequent accuracy gain [5], a data-centric lesson that recurs almost verbatim in Booth et al.'s MAP-Mapper, which reached up to 95% precision on MARIDA-derived data mainly by refining negative examples [13].

More recently, this data-hunger has become the target of new work. Dalsasso et al.'s CSM-Debris addresses it directly with a self-supervised, cross-sensor pretraining scheme: a shared U-Net-style encoder-decoder with sensor-specific input projection layers is pretrained to denoise both Sentinel-2 and PlanetScope imagery without any labels, then fine-tuned for debris segmentation only on labelled Sentinel-2 data. The resulting model was transferred to PlanetScope imagery that was never given labels for, reporting roughly a 20% F1/IoU improvement over Sentinel-2-only training with no loss on Sentinel-2 itself, and outperforming a general-purpose foundation model (SMARTIES) that lacked PlanetScope exposure during pretraining [24]. This is a plausible route toward the labelled-data scarcity that Rußwurm et al. and Booth et al. treat as the binding constraint, although it remains RGB-only, does not model cloud cover, and transfers only from Sentinel-2 to PlanetScope rather than to other sensors [24]. A parallel strand instead targets underwater, rather than floating-surface, debris: a lightweight CNN model referred to as UTNet is trained on a combined UTD2 and J-EDI underwater debris dataset and benchmarked against Faster R-CNN, SSD, and several YOLO variants, reporting higher detection accuracy alongside faster, more consistent real-time performance across different underwater depths, while, consistent with the rest of this literature, still degrading when water clarity is poor [25].

C. YOLO-Family Object Detection

Real-time UAV and USV deployments favor single-shot detectors. A2ANet builds on a YOLOv5s backbone with atrous convolutions and a channel-attention module, raising mAP@0.5 from 0.787 to 0.841 on a river-debris dataset at approximately 39 ms per image [8]. Zhou et al.'s YOLOTrashCan instead prioritised model size, shedding 30 MB via a feature-fusion network while addressing small-object and occlusion cases on the TrashCan-Instance benchmark [12]. Gains reported on either benchmark do not obviously transfer across rivers or lighting without retraining, a caveat that both papers acknowledge only briefly [8][12].

D. Hyperspectral Classification

El Bergui et al.'s lightweight spatial-spectral CNN (LSS-HCNN) reports above 97% mean accuracy across several hyperspectral benchmarks while cutting the parameter count by over 80% relative to a standard 2D-CNN [10], a substantial efficiency gain that remains validated mostly under controlled conditions. Alboody et al. demonstrated on-water deployment by embedding a comparable hyperspectral system on an unmanned aquatic drone [14]. Between the two, the on-water study is arguably the more operationally convincing evidence [14], even though its reported numbers are lower than those of the controlled-condition study[10].

E. SAR-Based and Multi-Modal Screening

Because optical sensors are blind under clouds and at night, SAR provides an all-weather complement. Pretorius et al. fused Sentinel-2 and RapidEye optical indices with Sentinel-1 backscatter over the South African EEZ, finding that VV polarization was more sensitive than VH, but noting that ships remained a persistent confounder across every channel and that no independent ground truth was available for validation [11]. The last caveat is worth flagging on its own: SAR fusion is consistently framed in this literature as promising rather than proven.

F. Trajectory Prediction and Drift Modelling

DEEP-PLAST couples YOLOv5 and U-Net++ detection with a classical Lagrangian advection model driven by Copernicus current and NOAA wind fields; however, its authors concede that drift validation remains largely qualitative given the scarcity of in-situ ground truth [9]. DriftNet takes a different route entirely, learning to reproduce drifter trajectories directly from sea-surface current fields via a Fokker–Planck-inspired architecture and reports a reduced mean separation distance relative to LSTM baselines [15]. Whether a learned or physics-based drift model generalizes better in data-sparse regions, such as much of the Indian Ocean, is still an open question that neither study answers directly [9][15].

G. Indian-Authored and India-Focused Research

Raju et al. combined a field survey with AI-based object detection along the central east coast of India and directly compared manual and automated counts on the same stretch of coastline [16]. Nivedita et al. applied FDI and Naïve-Bayes classification to Indian coastal waters, reporting 87.25% accuracy while explicitly targeting the sub-pixel detection problem [17], a figure close to Biermann's original 86% despite the different coastline, which is at least suggestive that FDI-plus-classifier performance is more stable across geographies than the Mediterranean-only literature might imply. Begum et al. used SVM and random forests over Nagapattinam, Tamil Nadu, and identified the red-edge, NIR, and SWIR-1 bands as the most discriminative for Indian coastal conditions [18]. Thanabalan et al. combined satellite sensors with drone imagery over Chennai, applying an HR-Net segmentation model across eight litter classes [22]. Field surveys from the Andaman and Nicobar Islands [19] and Car Nicobar [20] provide a ground-truth context that any future Indian Ocean satellite pipeline would eventually need for validation, and Sivadas et al.'s national review concludes that Indian marine litter research still lacks a coordinated, multidisciplinary roadmap [21]. Pattiaratchi et al.'s basin-scale synthesis situates this work within the wider Indian Ocean, noting that the region receives plastic from some of the most polluted rivers on Earth while remaining comparatively under sampled [23].

As Table I summarizes, the seven families trade off along broadly the same axes: spectral-index and classical ML methods are cheap and interpretable but threshold-sensitive; CNN/U-Net segmentation captures diffuse sub-pixel debris at the cost of data hunger; YOLO detectors are fast and edge-deployable but weaker on very small or occluded targets; hyperspectral classifiers offer the finest material discrimination at the highest sensor cost; SAR fusion adds all-weather coverage without standing alone as a plastic identifier; and Lagrangian/learned drift models project accumulation zones forward in time but remain, across the board, thinly validated. Table II summarizes representative studies to show how each family has been used, rather than how it performs on a benchmark alone.

Table I: Comparison Of Marine Debris Detection Method Families

Family	Principle	Reported Performance	Sensor Cost	Cofounder Robustness	Operational Maturity	Key References
Spectral index + thresholding	Thresholding of FDI/PI computed from multispectral imagery.	86–98% (site-dependent)	Low (Sentinel-2, free)	Weak vs. algae/foam	Field-tested, widely used	[1][2][6][7][17]
Classical ML (RF/SVM/SVR)	Supervised classifier on spectral-index features	87–98%	Low	Moderate	Field-tested	[6][17][19][20]
CNN / U-Net segmentation	Learned pixel-wise segmentation on multispectral patches	F1 0.50–0.95 (benchmark-dependent)	Low–Medium	Sensitive to label noise	Research-stage to early operational	[3]–[5][14]
YOLO-family detection	Single-shot object detector on UAV/USV imagery	mAP@0.5 0.79–0.84	Medium (UAV/USV)	Weak on occluded/small debris	Field-tested, edge-deployed	[9][13]
Hyperspectral CNN	Depthwise/pointwise CNN on hyperspectral cubes	>97% (controlled)	High	Strong (polymer-level)	Research-stage	[11][15]
SAR / multi-modal fusion	Backscatter anomaly detection, fused with optical	Qualitative (no ground truth)	Medium–High	Cloud-independent; ship confusion	Research-stage	[11]
Lagrangian/learned drift	Particle advection or learned surface-current model	Reduced separation vs. LSTM	Low (reanalysis data)	Largely unvalidated in situ	Research-stage	[10][16]

Table II: Summary Of Representative Literature

Study	Method Family	Region/Sensor	Key Findings
Biermann et al. [1]	FDI + Naive Bayes	Mediterranean, Sentinel-2	First operational spectral index; ~86% accuracy
Kikaki et al. [3]	RF vs. U-Net (MARIDA)	Global, Sentinel-2	F1 0.70 (RF) vs. 0.50 (U-Net); largest open benchmark
Sannigrahi et al. [6]	RF + SVM + indices	4 Mediterranean sites	92–98% accuracy; automated detection system
A2ANet [8]	YOLOv5s + atrous conv.	River/coastal USV	mAP@0.5 = 0.841; 39 ms/image, edge-deployable
DEEP-PLAST [9]	YOLOv5 + U-Net++ + drift	Black Sea, Sentinel-2/UAV	F1 = 0.840; drift validation qualitative only
El Bergui et al. [10]	Hyperspectral CNN (LSS-HCNN)	Hyperspectral cubes	>97% accuracy, 80% fewer parameters
Pretorius et al. [11]	Optical + SAR fusion	South African EEZ	No ground truth; ships are a persistent confounder
Raju et al. [16]	Field survey + AI detection	East coast of India	Direct manual-vs-AI comparison, India-specific
Nivedita et al. [17]	FDI + Naive Bayes	Indian coastal waters	87.25% accuracy; sub-pixel plastic/seaweed separation
Sivadas et al. [21]	Literature review	India (national)	Concludes that a coordinated research roadmap is still missing

IV. RESEARCH GAP

Fig. 2 illustrates the gap conceptually: detection-accuracy research and operational, agency-facing deployment literature overlap only partially, and the overlap rarely includes validated drift forecasting, regionally diverse benchmarking, or transparent reporting of what is live versus simulated data.

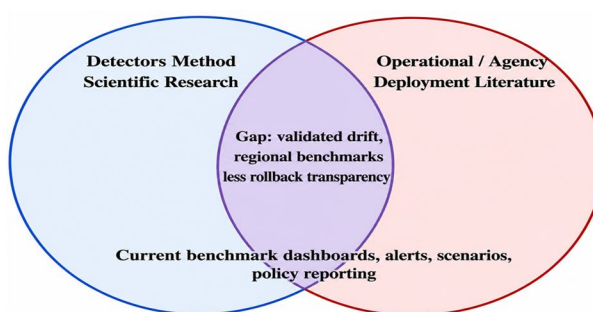


Figure 2 Conceptual illustration of the research gap between marine debris detection-accuracy research and operational deployment literature.

Five shortcomings are discussed in Sections II–III. Sub-pixel and biofouled-debris detection: no reviewed method reliably resolves debris below approximately 10–20% of a Sentinel-2 pixel's area or robustly compensates for algal biofouling on debris more than a few weeks old [1][17]. Fragmented regional benchmarking: Indian coastal studies [16][18][22] each rely on bespoke, single-site datasets; no shared, densely labelled, multi-region Indian benchmark comparable to MARIDA exists [3]. Unvalidated drift forecasting: the coupling of detection with drift modelling remains largely qualitative, even in the most integrated pipelines [9], and learned drift models have not been paired with detection-derived release points under Indian Ocean conditions [15]. Limited multi-sensor fusion at scale: individual studies have fused optical data with SAR [11] or with higher-resolution imagery [7], but none have integrated optical, SAR, and hyperspectral streams end-to-end for Indian coastal monitoring.

Deployment transparency: Much of the literature reports benchmark-grade accuracy without disclosing which pipeline stages are real, synthetic, or fallback — a gap the case study in Section V was explicitly built to address, although, as Section V-D makes clear, not every gap listed here is resolved by doing so.

V. CASE STUDY: AN OPERATIONAL DECISION-SUPPORT PIPELINE FOR INDIAN COASTAL WATERS

Building on the comparative analysis of Section III and the research gap of Section IV, this section presents the Marine Plastic Tracker, an implemented pipeline targeting INCOIS, the National Centre for Coastal Research (NCCR), and the Indian Coast Guard. The system is described honestly: several of its stages are genuinely live, while others currently run on synthetic or fallback data, and this distinction is stated explicitly rather than implied. No externally validated field-accuracy figure is claimed for any stage that has not been checked against an independent ground truth.

A. Pipeline Stages and Workflow

Fig. 3 details six timed, individually inspectable pipeline stages, each grounded in a concept from Section III. Download performs Sentinel-2 scene discovery via the Element84/AWS Open Data STAC API (live, no key required), with pixel-level ingestion available through the Copernicus Data Space Ecosystem's Sentinel Hub Process API. FDI computation applies the Biermann index as introduced in [1]. CNN classification uses a pluggable ONNX classifier that auto-detects patch-based ($64 \times 64 \times 6$) versus per-pixel input, falling back to an interpretable heuristic when no trained model is loaded; the shipped baselines are trained on synthetic spectra rather than MARIDA/MADOS [3] — disclosed openly rather than glossed over. Self-supervised, cross-sensor pretraining [24] offers a promising route to reduce this dependency and is noted as future work in Section VI. DBSCAN clustering groups detections at $\text{eps} = 0.09^\circ$ (~ 10 km) into HIGH/MEDIUM/LOW hotspots, echoing the broader detection-to-accumulation-zone logic used in prior detection-plus-drift pipelines [9] detections Drift forecasting produces a 72-hour Lagrangian ensemble with an uncertainty cone and estimated time-to-coast[15], though the underlying current/wind fields remain synthetic rather than live CMEMS/ERA5 data — the same validation gap flagged in Section III-F. Alerts aggregate HIGH/MEDIUM hotspots into a dashboard count.

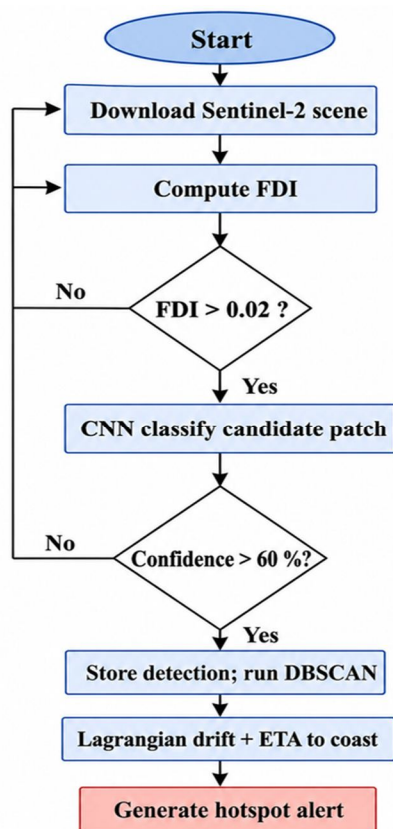


Figure 3 Workflow of the detection pipeline from scene download through hotspot alert generation.

B. Live vs. Simulated Components and Positioning Relative to the Literature

Every layer of the system degrades gracefully rather than failing outright, and Fig. 4 illustrates how deployment constraints, rather than convention, drive the selected family combination. Where an all-weather, cloud-penetrating requirement dominates, SAR fusion is the more defensible choice [11]; where polymer-level discrimination is required, hyperspectral classification is preferable [10]; however, for the deployment envelope actually targeted — a zero-budget, optical-only pipeline feeding a government-facing dashboard on a roughly five-day revisit cycle — the literature instead favors the FDI-plus-CNN-plus-clustering-plus-drift combination used here, on the strength of its low sensor cost, its field-tested accuracy across both Mediterranean and Indian coastal studies [1][6], [17], and its precedent in integrated pipelines such as [9].

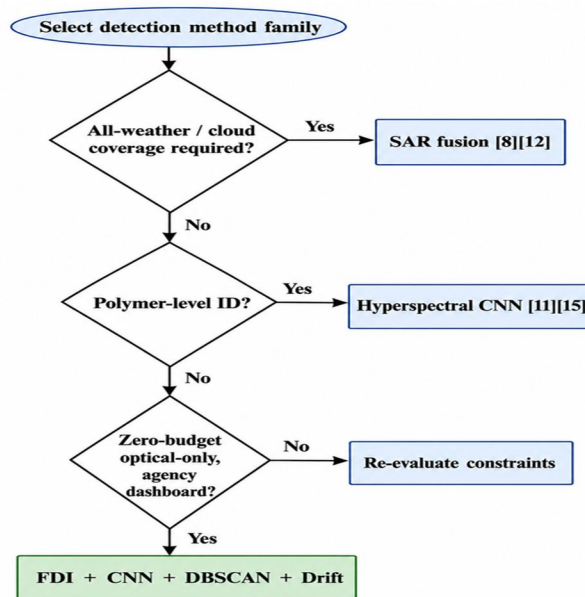


Figure 4 Decision flow showing how deployment constraints lead to the literature-supported family combination used in Marine Plastic Tracker.

C. Relation to the Identified Research Gap

The system addresses the deployment-transparency gap of Section IV directly through the per-stage live/demo tagging described above rather than through convention. It partially addresses the drift-validation gap by reporting an uncertainty cone rather than a single deterministic path, although the underlying current and wind fields remain synthetic; therefore, this should be read as a design choice that anticipates validation rather than validation itself. The sub-pixel/biofouling, regional benchmarking, and multi-sensor fusion gaps are not resolved by this system; they are noted in Table III as explicit future work, consistent with the honesty requirement this paper sets for itself.

TABLE III: RESEARCH GAP ANALYSIS

Identified Gap	Description	Addressed by Marine Plastic Tracker
Sub- pixel / biofouling	No method fully resolves sub-20% pixel debris or biofouled signal decay	Not resolved; documented as a known limitation
Regional benchmarking	No shared, multi-site Indian debris dataset exists	Not resolved; future work
Drift validation	Drift forecasts are rarely validated against in-situ tracks	Partially — uncertainty cone reported, but current/wind fields are still synthetic
Multi-sensor fusion	No end-to-end optical + SAR + hyperspectral pipeline	Not addressed; optical-only by design
Deployment transparency	Live vs. fallback status is rarely disclosed	Directly addressed via per-stage live/demo tagging

VI. LIMITATIONS AND THREATS TO VALIDITY

Two classes of limitations apply. First, the comparison in Section III is not a systematic review in the PRISMA sense, as no fixed search string, database set, or inclusion/exclusion protocol was pre-registered prior to the review. Coverage, while spanning 2020–2026 and 24 primary sources, may therefore not be exhaustive, and the selection of representative studies within each method family involved qualitative judgement rather than a reproducible screening procedure. Second, the cross-family comparison in Table I is itself constrained by how unevenly the underlying studies report their results: accuracy, F1-score, mAP, and qualitative assessments are not directly commensurable, and reported figures come from different sensors, study sites, and validation protocols rather than a shared benchmark. This makes any ranking across families indicative rather than statistically conclusive. Compounding this, the reviewed literature is itself imbalanced — global and Mediterranean studies dominate, drift-forecasting work is sparse and rarely validated against in-situ tracks, and Indian coastal studies are few, single-site, and not designed for cross-comparison. Consequently, the conclusions of this study should be read as literature-grounded design guidance, not as a statistically validated ranking of detection methods or a claim that the reviewed evidence base is complete.

VII. CONCLUSION

This paper argued that detection methods should be chosen not for how well they perform in isolation, but for how well they fit a given deployment. Toward that end, seven AI-driven detection-method families — spectral-index screening, classical machine learning, CNN/U-Net segmentation, YOLO-family detection, hyperspectral classification, SAR-based screening, and Lagrangian/learned drift forecasting — were compared on the terms that matter for continuous, agency-facing coastal monitoring rather than benchmark performance alone. This comparison surfaced a consistent pattern: drift forecasts go largely unvalidated against real tracks, Indian coastal studies remain fragmented and difficult to benchmark against one another, and deployment status — what is actually live versus simulated — is rarely disclosed with any precision.

Building on this synthesis, an implemented case-study pipeline was presented: a low-cost architecture combining FDI screening, CNN classification, DBSCAN clustering, and drift forecasting, purpose-built for Indian coastal waters. Its value lies not in a new detection algorithm, but in demonstrating what a literature-justified, transparently scoped architecture looks like when every stage is honestly labelled as live or fallback — a level of disclosure this review found conspicuously absent elsewhere.

This conclusion remains deliberately bounded. No externally validated field accuracy, drift accuracy, or latency figures are claimed for any stage of the pipeline. Its contribution is architectural and operational — assembling reviewed techniques into a defensible, transparently reported design — not algorithmic novelty. What this work does establish is a template: a structured basis for choosing between detection families, and a working example of the deployment honesty the field needs more of.

The path forward is concrete. It includes training the CNN stage on MARIDA/MADOS rather than synthetic spectra — potentially via self-supervised, cross-sensor pre-training approaches such as CSM-Debris that reduce dependence on site-specific labelled data [24] — wiring in live CMEMS/ERA5 drift fields, and reporting same-conditions field validation against Indian coastal ground truth. Beyond this specific pipeline, the wider field still needs what this review found missing throughout: a coordinated regional benchmark and genuine multi-sensor fusion, so that the next detection system built for India's coastline does not have to start, once again, from a single sentence justifying its sensor choice.

REFERENCES

- [1] L. Biermann, D. Clewley, V. Martinez-Vicente, and K. Topouzelis, "Finding plastic patches in coastal waters using optical satellite data," *Scientific Reports* vol. 10, p. 5364, 2020.
- [2] K. Themistocleous, C. Papoutsas, S. Michaelides, and D. Hadjimitsis, "Investigating detection of floating plastic litter from space using Sentinel-2 imagery," *Remote Sensing*, vol. 12, no. 16, p. 2648, 2020.
- [3] K. Kikaki, I. Kakogeorgiou, P. Mikeli, D. E. Raitsos, and K. Karantzas, "MARIDA: A benchmark for marine debris detection from Sentinel-2 remote sensing data," *PLOS ONE*, vol. 17, no. 1, p. e0262247, 2022.
- [4] J. Mifdal, N. Longépé, and M. Rußwurm, "Towards detecting floating objects on a global scale with learned spatial features using Sentinel-2," *ISPRS Annals of the Photogrammetry, Remote Sensing and Spatial Information Sciences*, vol. 3, pp. 285–293, 2021.
- [5] M. Rußwurm, S. J. Venkatesa, and D. Tuia, "Large-scale detection of marine debris in coastal areas with Sentinel-2," *iScience*, vol. 26, p. 108402, 2023.
- [6] S. Sannigrahi, B. Basu, A. S. Basu, and F. Pilla, "Development of automated marine floating plastic detection system using Sentinel-2 imagery and machine learning models," *Marine Pollution Bulletin*, vol. 178, p. 113527, 2022.
- [7] M. Kremezi, V. Kristollari, V. Karathanassi, K. Topouzelis, P. Kolokoussis, N. Taggio, A. Aiello, G. Ceriola, E. Barbone, and P. Corradi, "Increasing the Sentinel-2 potential for marine plastic litter monitoring through image fusion techniques," *Marine Pollution Bulletin*, vol. 182, p. 113974, 2022.
- [8] B. Badams, U. U. Sheikh, N. A. Wahab, S. A. R. Abu Bakar, M. I. Masud, M. Khoulj, U. Waheed, Z. A. Arfeen, and K. M. Goh, "A2ANet: Real-time detection of floating marine debris using atrous convolution and channel attention," *Ecological Informatics*, vol. 94, p. 103620, 2026.

- [9] A. Cernian and M.-E. Iliuta, "Empowering sustainability through AI-driven monitoring: The DEEP-PLAST approach to marine plastic detection and trajectory prediction for the Black Sea," *Water*, vol. 17, no. 22, p. 3318, 2025.
- [10] A. El Bergui, A. Porebski, and N. Vandenbroucke, "A lightweight spatial and spectral CNN model for classifying floating marine plastic debris using hyperspectral images," *Marine Pollution Bulletin*, 216, p. 117965, 2025.
- [11] J. Pretorius, S. Haupt, and B. Sibolla, "Exploring the potential of remote sensing to detect marine plastic debris in the South African Ocean region," *ISPRS Annals of the Photogrammetry, Remote Sensing and Spatial Information Sciences*, vol. X-G-2025, pp. 665–671, 2025.
- [12] W. Zhou, F. Zheng, G. Yin, Y. Pang, and J. Yi, "YOLOTrashCan: A deep learning marine debris detection network," *IEEE Transactions on Instrumentation and Measurement*, vol. 72, pp. 1–12, 2023.
- [13] H. Booth, W. Ma, and O. Karakuş, "High-precision density mapping of marine debris and floating plastics via satellite imagery," *Scientific Reports*, vol. 13, p. 6822, 2023.
- [14] A. Alboody, N. Vandenbroucke, A. Porebski, R. Sawan, F. Viudes, P. Doyen, and R. Amara, "A new remote hyperspectral imaging system embedded on an unmanned aquatic drone for the detection and identification of floating plastic litter using machine learning," *Remote*, 15, p. 3455, 2023.
- [15] O. Botvynko, C. Granero-Belinchon, S. G. Roux, N. B. Garnier, and A. Benzinou, "Deep learning for Lagrangian drift simulation at the sea surface," *arXiv preprint arXiv:2211.09818*, 2022; extended as "Neural prediction of Lagrangian drift trajectories on the sea surface," *Artificial Intelligence for the Earth Systems*, vol. 4, no. 3, 2025.
- [16] M. P. Raju, S. Veerasingham, V. Suneel, F. S. Asim, H. A. Khalil, M. Chatting, P. Suneetha, and P. Vethamony, "A machine learning-based detection, classification, and quantification of marine litter along the central east coast of India," *Frontiers in Marine Science*, vol. 12, 2025.
- [17] V. Nivedita, S. Sabarunisha Begum, G. Aldehim, A. M. Alashjaee, M. A. Arasi, M. Y. Sikkandar, T. Jayasankar, and S. Vivek, "Plastic debris detection along coastal waters using Sentinel-2 satellite data and machine learning techniques," *Marine Pollution Bulletin*, vol. 209, part A, p. 117106, 2024.
- [18] S. Sabarunisha Begum, K. Nithya, M. Sujatha, T. Jayasankar, N. B. Prakash, S. Srinivasan, and S. Vivek, "Automated system for identifying marine floating plastics to enhance sustainability in coastal environments through Sentinel-2 imagery and machine learning models," *Ocean Science Journal*, 2024.
- [19] S. Krishnakumar, S. Anbalagan, K. Kasilingam, P. Smrithi, S. Anbazhagi, and S. Srinivasalu, "Assessment of plastic debris in remote islands of the Andaman and Nicobar Archipelago, India," *Marine Pollution Bulletin*, vol. 151, p. 110841, 2020.
- [20] R. Kiruba-Sankar, K. Saravanan, S. Adamala, K. Selvam, K. L. Kumar, and J. Praveenraj, "First report of marine debris in Car Nicobar, a remote oceanic Island in the Nicobar archipelago, Bay of Bengal," *Regional Studies in Marine Science*, vol. 61, p. 102845, 2023.
- [21] S. K. Sivasdas, P. Mishra, T. Kaviarasan, M. Sambandam, K. Dhineka, M. V. R. Murthy, S. Nayak, D. Sivyer, and D. Hoehn, "Litter and plastic monitoring in the Indian marine environment: A review of current research, policies, waste management, and a roadmap for multidisciplinary action," *Marine Pollution Bulletin*, vol. 176, p. 113424, 2022.
- [22] P. Thanabalan, K. Gayathri, P. Mishra, T. Usha, S. K. Dash, and S. R. Marigoudar, "Monitoring of marine floating debris and beach litter using satellite and drones: A synergistic approach on policy and decision making," *Journal of the Indian Society of Remote Sensing*, 2025.
- [23] C. Pattiaratchi, M. van der Mheen, C. Schlundt, B. E. Narayanaswamy, A. Sura, S. Hajbane, R. White, N. Kumar, M. Fernandes, and S. Wijeratne, "Plastics in the Indian Ocean – sources, transport, distribution, and impacts," *Ocean Science*, vol. 18, pp. 1–28, 2022.
- [24] E. Dalsasso, M. Rußwurm, C. Donner, S. Darmon, R. de Vries, M. Volpi, and D. Tuia, "Self-supervised pre-training enables marine debris detection across sensors," *Remote Sensing of Environment*, vol. 339, 2026.
- [25] J. Cui, S. Zhou, G. Xu, X. Liu, and X. Gao, "Marine debris detection in real time: A lightweight UINet model," *Journal of Marine Science and Engineering*, vol. 13, p. 1560, 2025.



10.22214/IJRASET



45.98



IMPACT FACTOR:
7.129



IMPACT FACTOR:
7.429



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Call : 08813907089  (24*7 Support on Whatsapp)