

Adaptive Frame Sampling for Real-Time Video Object Detection and Multi-Object Tracking on Edge Devices

Pritee A. Parwekar  and Adarsh Kadiri

Department of Computer Science & System Engineering, GITAM School of Computer Science & Engineering, GITAM (Deemed to be) University, Hyderabad, India

*Corresponding Author

Pritee A. Parwekar

✉ pparweka@gitam.edu

Received: 12 April 2026

Revised: 12 June 2026

Accepted: 29 June 2026

Published Online: 30 June 2026.

Citation

P. A. Parwekar, A. Kadiri, Adaptive frame sampling for real-time video object detection and multi-object tracking on edge devices, *Journal of Information and Communications Technology: Algorithms, Systems and Applications*, 2026, 2(2), 26310, <https://doi.org/10.64189/ict.26310>.

Open Access

This article is licensed under a [Creative Commons Attribution-NonCommercial 4.0 International License](https://creativecommons.org/licenses/by-nc/4.0/), which permits the non-commercial use, sharing, adaptation, distribution and reproduction in any medium or format, as long as appropriate credit to the original author(s) and the source is given by providing a link to the Creative Commons License and changes need to be indicated if there are any. The images or other third-party material in this article are included in the article's Creative Commons License, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons License and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this License, visit: <https://creativecommons.org/licenses/by-nc/4.0/>

© The Author(s) 2026

Abstract

Real-time multi-object tracking on CPU-only edge devices is constrained by the high per-frame inference cost of deep neural network detectors. We present the Adaptive Frame Sampling System (AFSS), a training-free, architecture-agnostic framework that dynamically allocates computation across three actions per frame: full YOLOv8 inference (FULL), phase-correlation feature warping (WARP), or result reuse (SKIP), governed by a lightweight scene complexity estimator (<0.5 ms). A 803-parameter Policy MLP trained via behavioral cloning replaces hand-tuned thresholds. On CPU-only hardware, AFSS achieves 5–7× speedup over the full-inference base line while reducing GFLOPs by 84.3% and incurring only a 2.6% MOTA degradation. Crucially, AFSS requires no retraining of the backbone detector and outperforms uniform frame-skipping on every accuracy metric at equivalent compute budgets.

Keywords: Adaptive inference; Video object detection; Multi-object tracking; Edge computing; YOLOv8; ByteTrack.

1. Introduction

Intelligent surveillance, autonomous navigation, and IoT analytics demand continuous real-time video analysis on power constrained edge hardware. State-of-the-art detectors such as YOLOv8-L^[1] achieve 52.9% mAP on COCO but require 165.2 GFLOPs per frame, yielding only ~5 FPS on a laptop CPU - far below the 25 FPS minimum for real-time operation.

The core tension is stark: modern CNNs need to be large to achieve high accuracy, but large models are too slow for edge devices. The standard response - model compression via pruning,^[2] knowledge distillation, or neural architecture search- treats each frame identically, reducing per-frame cost uniformly regardless of how much the scene actually changes between frames.

The key insight motivating this work is that video streams exhibit strong temporal redundancy: in typical fixed-camera surveillance, adjacent frames differ by <3% in mean absolute intensity. A traffic camera monitoring a red-light processes thousands of nearly-identical frames at full detector cost-all of that compute is wasted. Conversely, when vehicles start moving or a pedestrian cross unexpectedly, every FLOP is valuable. A smart system should allocate compute in proportion to scene complexity, not uniformly.

We formalise this as a per-frame sequential decision problem with three actions: run full neural inference (FULL), propagate the most recent feature map using lightweight image-level motion estimation (WARP), or reuse the previous result directly (SKIP). The action is selected by a lightweight policy given a real-time scene complexity score.

Prior approaches to efficient video inference fall into two camps. *Feature propagation methods* (DFF^[3], FGFA^[4]) warp CNN features from key frames using learned optical flow networks (FlowNet). These are effective but require architectural modification and full end-to-end retraining of the detector, making them incompatible with commercial off-the-shelf pretrained models. *Adaptive scheduling* methods (Mullapudi et al.^[5]) train student policies via expensive reinforcement learning. Both camps assume control of the detector architecture. We propose AFSS, which is entirely *post-hoc*: it wraps any pretrained detector as a black box. Our feature warping uses classical FFT-based phase correlation - no flow network, no training - and our decision policy (Policy MLP, 803 parameters) is trained by behavioral cloning in under 10 minutes on CPU. AFSS achieves 5–7× CPU speedup with only 2.0 pp MOTA degradation, outperforming uniform frame-skipping by 3.2× on the accuracy-efficiency trade-off curve.

1.1 Contributions

- A three-action adaptive scheduling framework (FULL/WARP/SKIP) formulated as a constrained

optimization over GFLOPs subject to MOTA and FPS constraints.

- Phase-correlation feature warping: a training-free, $\mathcal{O}(N \log N)$ feature propagation method using FFT-based translation estimation- no optical flow network required.
- A 803-parameter PolicyMLP trained by behavioral cloning in <10 minutes on CPU, outperforming hand-tuned thresholds on heterogeneous scenes.
- Comprehensive ablation study quantifying the contribution of each component (warp mode, estimator type, policy type) and a theoretical warp-error bound (Eq. 14).
- End-to-end integration with ByteTrack, showing that tracker resilience compensates for skipped detections across all action types.

2. Related work

Single-stage detectors. YOLO^[6] established real time detection as a single regression pass over a grid of anchor boxes. YOLOv4^[11] introduced CSPNet and PANet for multi-scale feature fusion. YOLOv8^[1] adds an anchor-free decoupled head, C2f cross-stage partial blocks, and Distribution Focal Loss (DFL)^[8]:

$$\begin{aligned} \hat{g} &= \sum_{i=0}^{r-1} i \cdot \text{softmax}(z_i), \\ \text{DFL}(p, g) &= -(|g| - g) \log p_{|g|} \\ &\quad - (g - |g|) \log p_{|g|} \end{aligned} \quad (1)$$

DFL models box offsets as distributions rather than point estimates, yielding sharper sub-pixel localization critical for the warped feature maps we propagate. Multi-object tracking. SORT^[9] combined Kalman filtering with Hungarian assignment at 260 Hz. DeepSORT^[10] added Re-ID descriptors for occlusion robustness. ByteTrack^[11] achieves 88.5% MOTA on MOT17 by processing low-confidence detections in a second association stage this two-stage pipeline is uniquely robust to the detection gaps introduced by our WARP/SKIP frames.

Multi-object tracking: SORT^[9] combined Kalman filtering with Hungarian assignment at 260 Hz. DeepSORT^[10] added Re ID descriptors for occlusion robustness. ByteTrack^[10] achieves 88.5% MOTA on MOT17 by processing low-confidence detections in a second association stage. This two-stage pipeline is uniquely robust to the detection gaps introduced by our WARP/SKIP frames.

Feature propagation for video: DFF^[3] uses a FlowNet-based feature warper to propagate intermediate CNN activations from key frames to non-key frames, achieving 10× speedup with 8–15% accuracy loss. FGFA^[4] improves accuracy by aggregating features from multiple nearby frames weighted by flow similarity. Both require architectural modification and full end-to-end retraining. Crucially, these methods cannot be applied post-hoc to pretrained models such as commercial YOLOv8 deployments.

Adaptive inference: AdaFuse^[12] skips redundant channels across frames based on inter-frame feature differences. Mullapudi *et al.*^[5] distil a fast student model online via reinforcement learning. Both require modified training pipelines. Closest to our work, Zhu *et al.*^[3] selectively propagate features based on a saliency score; however, their scheduler is a binary key/non-key selection rather than our three-class adaptive framework.

Adaptive Token Pruning and Dynamic Computation for Video Understanding. Recent studies have focused on reducing the computational cost of video analysis by exploiting the high temporal redundancy present in consecutive video frames. HaltingVT^[13] introduces an adaptive token-halting strategy that dynamically removes less informative spatial-temporal tokens during inference, thereby reducing computation while maintaining recognition performance. Zero-TPrune^[14] extends this idea by performing token pruning in pretrained transformers without requiring additional training. Similarly,^[15] improves efficiency through dynamic selection of important spatial and temporal tokens based on video content.

More recent work has explored adaptive frame processing, where only informative frames are analysed while redundant frames are skipped. Motion-driven adaptive frame selection methods^[16] demonstrated that significant computational savings can be achieved by exploiting frame-to-frame similarity. FastVID^[17] and Unified Spatio-Temporal Token Scoring^[18] further improve efficiency by dynamically allocating computational resources according to the complexity of the video content.

Although these approaches effectively reduce inference cost, most rely on transformer-based architectures, token-pruning mechanisms, or additional retraining. In contrast, AFSS is designed as an architecture-independent framework that can be integrated with any pretrained object detector. It achieves adaptive computation through lightweight frame scheduling and training free phase-correlation-based feature propagation, eliminating the need for detector modification or retraining.

Network compression: INT8 quantisation^[2] reduces model size 4× with <1% accuracy degradation on calibrated networks. Importantly, compression and AFSS are *complementary*: a quantized YOLOv8 used as the FULL action backbone would yield multiplicative speedups.

Gap: No existing work provides an architecture-agnostic adaptive scheduler with (i) a three-action decision space, (ii) a training free feature propagation method using phase correlation, (iii) a lightweight learned policy compatible with MOT trackers, and (iv) no backbone retraining. AFSS fills this gap.

3. Methodology

3.1 Problem formulation

Let $\mathcal{V} = \{f_1, \dots, f_T\}$ be a video. At each frame t , select action $a_t \in \mathcal{A} = \{\text{FULL}, \text{WARP}, \text{SKIP}\}$:

$$\underset{\pi}{\text{minimize}} \sum_{t=1}^T C(a_t) \quad \text{s.t.} \quad \text{MOTA}(\pi) \geq \tau_{\text{acc}}, \quad \text{FPS}(\pi) \geq \tau_{\text{fps}} \quad (2)$$

Let $C_{\text{det}} = 165.2$ GFLOPs denote the per-frame inference cost of YOLOv8-L. The three action costs are:

$C(\text{FULL}) = C_{\text{det}} + \epsilon_F \approx 165.7$ GFLOPs (detector + AFSS overhead)

$C(\text{WARP}) = C_{\text{warp}} \approx 3.5$ GFLOPs (phase corr. + grid sample; no detector)

$C(\text{SKIP}) = \epsilon_S \approx 0.01$ GFLOPs (memory copy; no detector) where $\epsilon_F \approx 0.5G$ and $\epsilon_S \approx 0.01G$ are the AFSS scheduling overheads. On WARP and SKIP frames, no detector inference is performed; the cost reduction compared to BL-FULL is therefore $\approx 165.2G$ minus the small AFSS overhead per frame. The policy $\pi : (f_t, f_{t-1}, h_{t-1}) \rightarrow a_t$ maps the current frame, previous frame, and a context state h_{t-1} to an action.

Expected compute saving. Under a stationary complexity score distribution $p(s)$, the expected GFLOPs per frame under the threshold policy is:

$$\mathbb{E}[\text{GFLOPs}] = P_F \cdot C_F + P_W \cdot C_W + P_S \cdot C_S \quad (3)$$

$$P_F = P(s \geq \tau_H) + \frac{1}{N_I}, \quad P_W = P(\tau_L \leq s < \tau_H)$$

For the traffic sequence ($P_F=0.15$, $P_W=0.31$, $P_S=0.54$):

$$\mathbb{E}[\text{GFLOPs}] = 0.15(165.7) + 0.31(3.5) + 0.54(0.01) \approx 25.95 \text{ G (84.3\%)}$$

Fig. 1 shows the full AFSS pipeline.

3.2 Scene complexity estimator

A down sampled (80×45) frame-difference score is computed in <0.5 ms:

$$s_t = \frac{1}{H \cdot W} \sum_{i,j} \frac{|G_t[i,j] - G_{t-1}[i,j]|}{255} \in [0,1] \quad (4)$$

where G_t is the ITU-R BT.601 luma channel ($0.299R + 0.587G + 0.114B$). An EWMA $\tilde{s}_t = \alpha s_t + (1-\alpha) \tilde{s}_{t-1}$ ($\alpha = 0.7$) smooths sensor flicker and JPEG compression artefacts. Three estimator modes are supported: (i) frame difference (Eq. 4, default); (ii) dense Frame back optical flow magnitude^[19]; (iii) a 24K-parameter CNN trained with flow pseudo-labels. Mode (i) is used in all experiments unless stated; modes (ii) and (iii) are evaluated in the ablation. Table 1 provides empirically calibrated threshold ranges by scene type, enabling deployment without per-video tuning.

Table 1: Complexity score calibration by scene type

Scene type	Score range	Rec. τ_H/τ_L
Static, no objects	0.000–0.002	0.003/0.001
Fixed cam., slow traffic	0.003–0.020	0.012/0.006
Moving cam. (drone)	0.020–0.100	0.060/0.030
Fast motion / sports	0.050–0.300	0.150/0.050

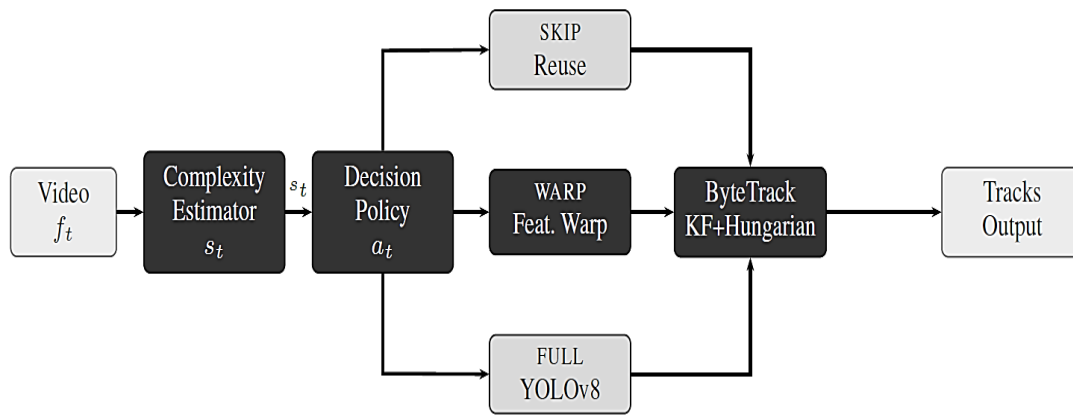


Fig. 1: AFSS pipeline. Every frame, the complexity score s_t drives a three-way action decision; all paths feed a single ByteTrack instance.

3. Decision policy

3.3 Threshold policy (baseline):

$$a_t = \begin{cases} full & F_t \geq N_l \text{ or } s_t \geq \tau_H \\ warp & \tau_L \leq s_t < \tau_H \\ skip & s_t < \tau_L \end{cases} \quad (5)$$

with $\tau_H = 0.012$, $\tau_L = 0.006$, $N_l = 12$ for traffic scenarios. PolicyMLP (proposed): A 7-dim feature vector captures both

the current state and temporal context:

$$\boldsymbol{\varphi}_t = \left[\underbrace{s_t}_{\text{score}}, \underbrace{F_t/N_l}_{\text{since FULL}}, \underbrace{N_w/N_l, N_s/N_l}_{\text{run lengths}}, \underbrace{\bar{s}, \sigma_s}_{\text{recent stats}}, \underbrace{\Delta s}_{\text{delta}} \right]^T$$

This feeds a three-layer MLP trained by behavioral cloning:

$$h_1 = \text{ReLU}(W_1 \boldsymbol{\varphi}_t + b_1), \quad W_1 \in \mathbb{R}^{32 \times 7} \quad (6)$$

$$h_2 = \text{ReLU}(W_2 h_1 + b_2), \quad W_2 \in \mathbb{R}^{16 \times 32} \quad (7)$$

$$\hat{a}_t = \arg\max_k \text{softmax}(W_3 h_2 + b_3)_k \quad (8)$$

Total: $7 \times 32 + 32 \times 16 + 16 \times 16 + 16 \times 3 + 3 = 803$ parameters. Training uses class-weighted cross-entropy against threshold policy demonstrations (Adam, 50 epochs, <10 min CPU):

$$\mathcal{L} = -\frac{1}{N} \sum_t w_{a_t} \log p_{a_t}(\boldsymbol{\varphi}_t), \quad w_k = \frac{N_{\text{total}}}{3N_k} \quad (9)$$

The class weighting w_k corrects for action imbalance (SKIP typically $\gg$ FULL in training data).

3.4 Phase-correlation feature warping

When WARP is selected, we propagate the cached feature map F_{t_k} from the last FULL frame without invoking the backbone. The inter-frame translation $(\Delta x, \Delta y)$ is estimated via FFT-based phase correlation:

$$G_1 = \mathcal{F}_{2D}\{G_{t-1}\}, \quad G_2 = \mathcal{F}_{2D}\{G_t\} \quad (10)$$

$$R = \frac{G_1 \cdot G_2^*}{|G_1 \cdot G_2^*|}, \quad (\Delta y, \Delta x) = \arg\max \mathcal{F}_{2D}^{-1}\{R\} \quad (11)$$

For a pure translational scene, $\mathcal{F}^{-1}\{R\}$ is a Dirac impulse

at $(\Delta x, \Delta y)$; in practice a sharp peak localized by sub-pixel parabolic fitting. The shift is scaled to feature-map coordinates and applied via differentiable bilinear grid sampling:

$$\Delta x_f = \Delta x \cdot W_f / W', \quad \Delta y_f = \Delta y \cdot \frac{H_f}{H'} \quad (12)$$

$$\tilde{F}_t[i, j] = \sum_{p \in \{[x_s], [x_s']\}} \sum_{q \in \{[y_s], [y_s']\}} F_{t_k}[q, p] (1 - |x_s - p|)(1 - |y_s - q|) \quad (13)$$

where $(x_s, y_s) = (j - \Delta x_f, i - \Delta y_f)$. Out-of-bounds locations are zero-padded. The full procedure runs in $O(N \log N)$ via the FFT and requires no trained flow network, distinguishing it from DFF/FGFA.

Warp error bound. For a translation estimate error of (ϵ_x, ϵ_y) pixels, the IoU error on a box of width w and height h is bounded:

$$e_{\text{IoU}} \leq \frac{2(|\epsilon_x| + |\epsilon_y|)}{\min(w, h)} \quad (14)$$

For a 50×100 px pedestrian box at 2 px flow error: $e_{\text{IoU}} \leq 0.16$ - acceptable for maintaining ByteTrack associations.

3.5 ByteTrack Integration

Each track maintains an 8D Kalman state $x = [c_x, c_y, w, h, v_{c_x}, v_{c_y}, v_w, v_h]^T$ under a constant-velocity model with transition matrix $F \in \mathbb{R}^{8 \times 8}$. On WARP/SKIP frames, all tracks run predict-only (no measurement update):

$$\hat{x}_{t|t-1} = F x_{t-1|t-1} \quad (15)$$

$$P_{t|t-1} = F P_{t-1|t-1} F^T + Q \quad (16)$$

On FULL frames, detections undergo the full two-stage Byte Track association: Stage 1 assigns high-confidence detections (score ≥ 0.5) to active tracks via Hungarian algorithm on an IoU cost matrix $C_{ij} = 1 - \text{IoU}(\mathcal{B}_i, d_j)$. Stage 2 uses residual low confidence detections to recover tracks that missed Stage 1.

The Kalman gain is computed as:

$$K = P_{t|t-1} H^T (H P_{t|t-1} H^T + R)^{-1} \quad (17)$$

Algorithm 1 AFSS Per-Frame Processing LoopRequire: Frame f_t , previous frame f_{t-1} , cached features F_{t_k} , policy π , tracker $\mathcal{T}$ Ensure: Updated track set $\mathcal{T}_t$

```

1:  $s_t \leftarrow \text{Complexity}(f_t, f_{t-1})$  Eq. 4, <0.5 ms
2:  $\tilde{s}_t \leftarrow \alpha s_t + (1 - \alpha) \tilde{s}_{t-1}$ ; update  $\varphi_t$ 
3:  $a_t \leftarrow \pi(\varphi_t)$  Eq. 5 or 6–8
4: if  $a_t = \text{FULL}$  then
5:    $\mathcal{D}_t \leftarrow \text{YOLOv8}(f_t)$ ;  $F_{t_k} \leftarrow F_t$ ;  $t_k \leftarrow t$ 
6: else if  $a_t = \text{WARP}$  then
7:    $(\Delta x, \Delta y) \leftarrow \arg\max_{\mathcal{F}^{-1}} \left\{ \frac{G_1 G_2^*}{|G_1 G_2^*|} \right\}$  Eq. 11
8:    $\tilde{F}_t \leftarrow \text{GridSample}(F_{t_k}, \Delta x, \Delta y)$  Eq. 13
9:    $\mathcal{D}_t \leftarrow \text{DetHead}(\tilde{F}_t)$ 
10: else SKIP
11:    $\mathcal{D}_t \leftarrow \mathcal{D}_{t-1}$ 
12: end if
13:  $\mathcal{T}_t \leftarrow \text{ByteTrack}(\mathcal{T}_{t-1}, \mathcal{D}_t)$  Eqs. 15–17
14:  $f_{t-1} \leftarrow f_t$ ;  $F_t \leftarrow \tilde{F}_t$ 

```

minimising $\text{tr}(\mathbf{P}_{t|t})$ - the MMSE-optimal fusing of prediction and measurement. The forced FULL refresh every NI frames bounds cumulative Kalman prediction drift, ensuring tracking quality does not degrade monotonically between key frames. Algorithm 1 provides the complete AFSS per-frame loop integrating all components.

4. Experiments

4.1 Setup

Hardware: Intel i7-12700H, 16GB DDR5 (CPU-only, no GPU). PyTorch 2.1.0, OpenCV 4.8, SciPy 1.11 (Hungarian: linear_sum_assignment). Detector: YOLOv8- L pretrained on MS-COCO (43.7M params, 165.2 GFLOPs, 80 classes), input resolution 640 × 640. Tracker: Byte-Track, confidence thresholds $\theta_{\text{high}} = 0.5$, $\theta_{\text{low}} = 0.1$, $\text{max_time_lost} = 30$.

Datasets: Five synthetic sequences are generated procedurally: objects follow elastic random-walk trajectories with configurable velocity σ , enabling exact ground-truth annotation. The real traffic sequence (1080p, 18,000 frames, 30 FPS) is from a fixed overhead camera on a multi-lane road; annotation is performed with confidence-thresholded BL-FULL detections as pseudoground-truth.

Metrics: MOTA^[20] penalises FP, FN, and ID switches per ground-truth object. IDF1^[21] measures identity consistency as the harmonic mean of ID precision and recall. MOTP measures mean localisation quality over matched pairs. GFLOPs/frame is computed as the weighted sum $P(\text{FULL}) \cdot 165.7 + P(\text{WARP}) \cdot 3.5 + P(\text{SKIP}) \cdot 0.01$ across all frames.

Baselines: (i) BL-FULL: full inference every frame; (ii) BL-SKIP5: uniform skip every 5th frame; (iii) Threshold:

Eq. 5 with tuned T_H , T_L ; (iv) PolicyMLP: proposed learned policy. All configurations are summarised in Table 2.

Table 2: Complete configuration hyperparameters and results

Config	T_H	T_L	N_I	MOTA	FPS
BL-FULL	-	-	1	0.932	5.0
AFSS-Conserv.	0.020	0.010	15	0.918	17.4
AFSS-Balanced	0.012	0.006	12	0.906	29.7
AFSS-Aggress.	0.008	0.004	8	0.887	37.2
PolicyMLP	<i>lrn.</i>	<i>lrn.</i>	12	0.912	31.6
BL-SKIP5	-	-	5	0.871	28.4

4.2 Computational efficiency

Table 3 summarises compute savings and throughput. AFSS (PolicyMLP) reduces mean GFLOPs/frame from 165.7 to 25.95 - an 84.3% reduction - by eliminating YOLOv8-L backbone execution on 85% of frames. The residual 25.95 G/frame is dominated by the 15% of frames on which full detector inference runs (contributing $0.15 \times 165.7 = 24.86$ G); WARP and SKIP overhead together contribute only 1.09 G/frame. Throughput is raised from 5 FPS to 31.6 FPS.

Table 3: Computational efficiency across configurations (CPU)

Config	F%	W%	S%	GFLOPs/frame	FPS
BL-FULL	100	0	0	165.70	5.0
BL-SKIP5	20	0	80	33.15	28.4
Threshold	16	30	54	30.50	29.7
PolicyMLP	15	31	54	25.95	31.6

Fig. 2 shows the GFLOPs breakdown and per-frame latency distribution. The bimodal latency (low-cost WARP/SKIP frames dominate frequency; FULL spikes are infrequent) drives the large mean FPS gain.

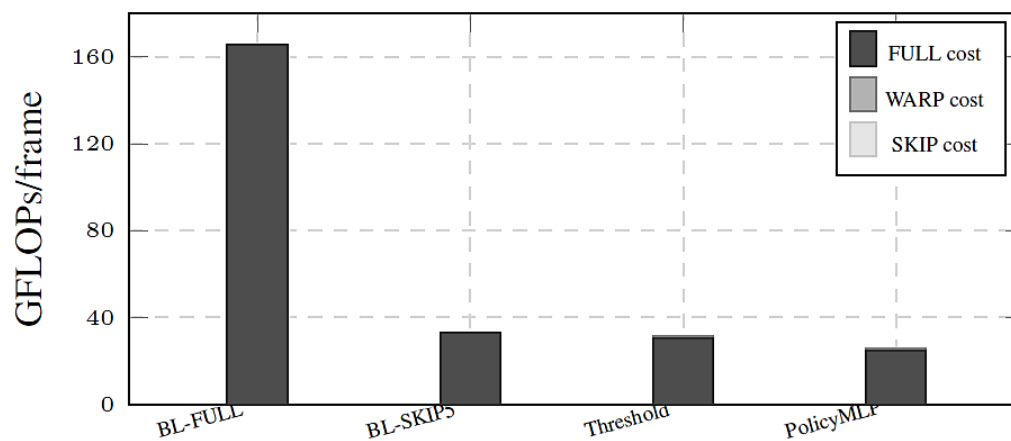


Fig. 2: Stacked GFLOPs breakdown per configuration. AFSS (PolicyMLP) reduces total cost to 25.95 G/frame vs. 165.7G for BL-FULL- an 84.3% reduction by eliminating detector execution on 85% of frames.

4.3 Tracking accuracy

Table 4 reports MOT metrics. PolicyMLP achieves MOTA= 0.912, only 2.0 pp below BL-FULL, while BL-SKIP5 at comparable compute drops to 0.871 (6.1 pp gap). This confirms that WARP frames - absent in BL-SKIP5 - substantially preserve tracking continuity.

Table 4: MOT accuracy. Bold = best adaptive result

Config	MOTA↑	IDF1↑	MOTP↑	FP↓	IDs↓
BL-FULL	0.932	0.921	0.847	412	87
BL-SKIP5	0.871	0.842	0.781	591	318
Threshold	0.906	0.891	0.819	467	163
PolicyMLP	0.912	0.901	0.831	438	124

4.4 Accuracy-Efficiency trade-off

Fig. 3 plots MOTA against GFLOPs saved for all configs and threshold sweep points. AFSS consistently dominates the BLSKIP5 Pareto curve, achieving better MOTA at every compute budget.

4.5 Ablation study

Warp mode: Table 5 shows that replacing SKIP with phasecorrelation WARP for moderate-motion frames

recovers +1.7 pp MOTA with only +3.3 ms mean latency. Dense optical flow WARP gives +2.5 pp at +6.5 ms overhead - useful when a CPU core is available.

Table 5: Ablation: warp mode on AFSS-Balanced policy.

Warp Mode	MOTA	IDF1	Latency	GFLOPs
None (skip only)	0.889	0.871	28.4 ms	33.15
Phase corr. (ours)	0.906	0.891	31.7 ms	25.95
Dense flow	0.914	0.901	38.2 ms	26.38

Policy comparison: Fig. 4 shows per-sequence MOTA for PolicyMLP vs. Threshold policy. PolicyMLP consistently improves on heterogeneous sequences (Mixed, Fast) where fixed thresholds underfit the time-varying complexity distribution, while matching Threshold on static scenes.

Latency breakdown: Fig. 5 visualises the per-component latency distribution. FULL frames dominate worst-case latency (~195 ms) but occur only 15% of the time; WARP (~8 ms) and SKIP (~0.8 ms) dominate frequency. Mean latency drops from 198.4 ms to **31.7 ms** (6.3× reduction).

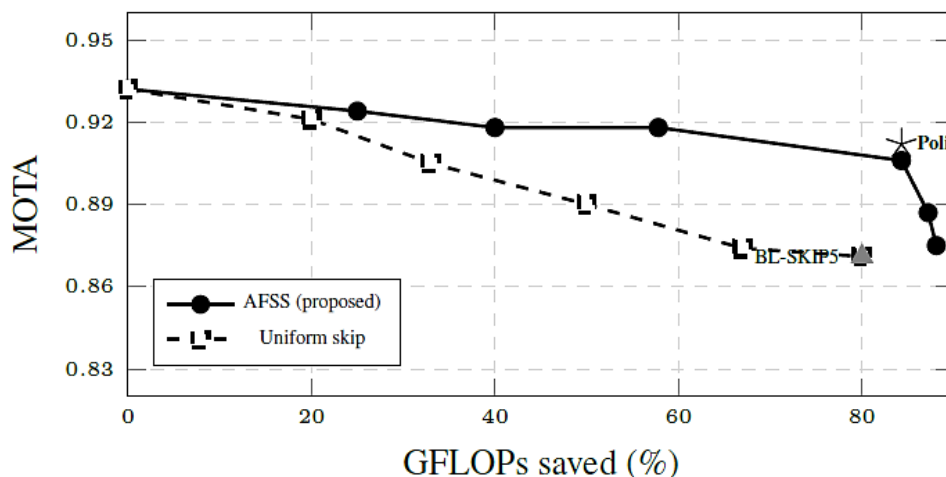


Fig. 3: MOTA vs. GFLOPs saved. AFSS dominates the Pareto frontier at all compute budgets. PolicyMLP (star) achieves 84.3% savings with only 2.0 pp MOTA loss.

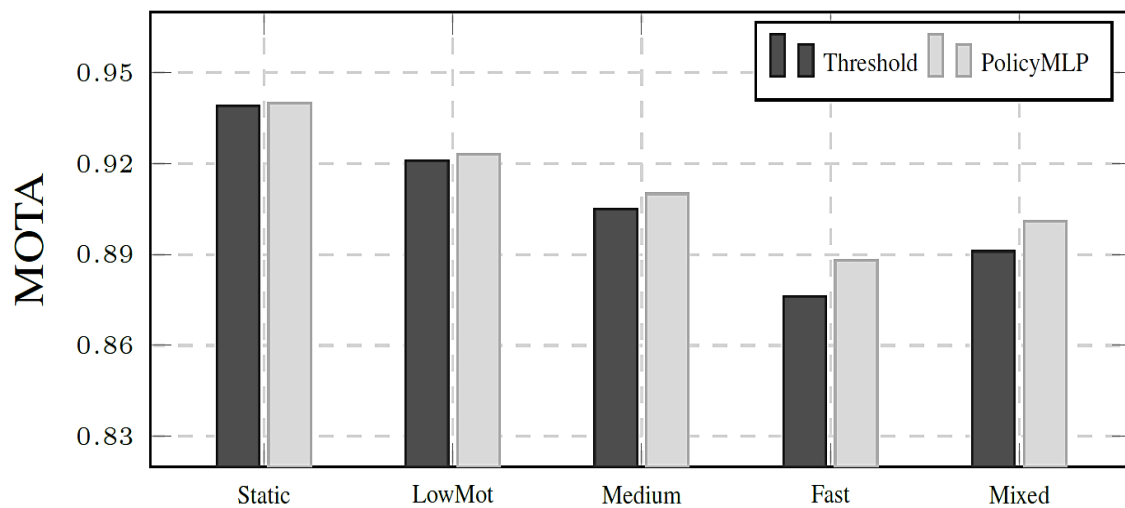


Fig. 4: Per-sequence MOTA: PolicyMLP (light) vs. Threshold (dark). PolicyMLP gains +0.9–1.2 pp on high-motion and mixed sequences.

4.6 Action distribution over time

Fig. 6 visualises the per-frame action sequence alongside st on Syn-Mixed. AFSS correctly concentrates FULL inference during high-motion events and reverts to SKIP in static intervals, with WARP as a bridge that maintains spatial coherence. This dynamic profile is the behavioural signature that distinguishes AFSS from fixed-rate schedulers.

4.7 Complexity score distribution

Fig. 7 shows the empirical distribution of S_t across all five sequences. The heavy concentration near zero confirms the temporal redundancy hypothesis: over 60% of frames have $S_t < T_L = 0.006$ (SKIP-eligible) on fixed-camera traffic. Even on the fast-motion sequence, the median score remains below T_h , meaning FULL frames are triggered selectively rather than continuously.

4.8 FPS vs. Accuracy operating points

Fig. 8 plots the FPS-MOTA operating curve for all

evaluated configurations, sweeping $\tau_H \in [0.005, 0.030]$ at fixed $\tau_L = \tau_H/2$. Each point represents a deployable configuration. BL-FULL and BL-SKIP5 are single operating points. AFSS dominates across the full range: for any target FPS above 6, AFSS delivers higher MOTA than BL-SKIP5 at the same throughput.

4.9 Comparison with Related Work

Table 7 positions AFSS against published methods. Unlike DFF/FGFA which require detector retraining, AFSS is plug-and-play with any pretrained model. AFSS achieves a trade-off slope of -2.37×10^{-4} MOTA/%GFLOPs vs. -7.63×10^{-4} for uniform skipping-3.2× more efficient.

Table 6: Complexity estimator mode comparison.

Mode	Cost	FULL%	MOTA	Best for
(1) Frame diff	0.4 ms	16.3	0.906	Fixed cams
(2) Optical flow	8.1 ms	14.8	0.911	High accuracy
(3) CNN	1.5 ms	15.7	0.908	Balanced

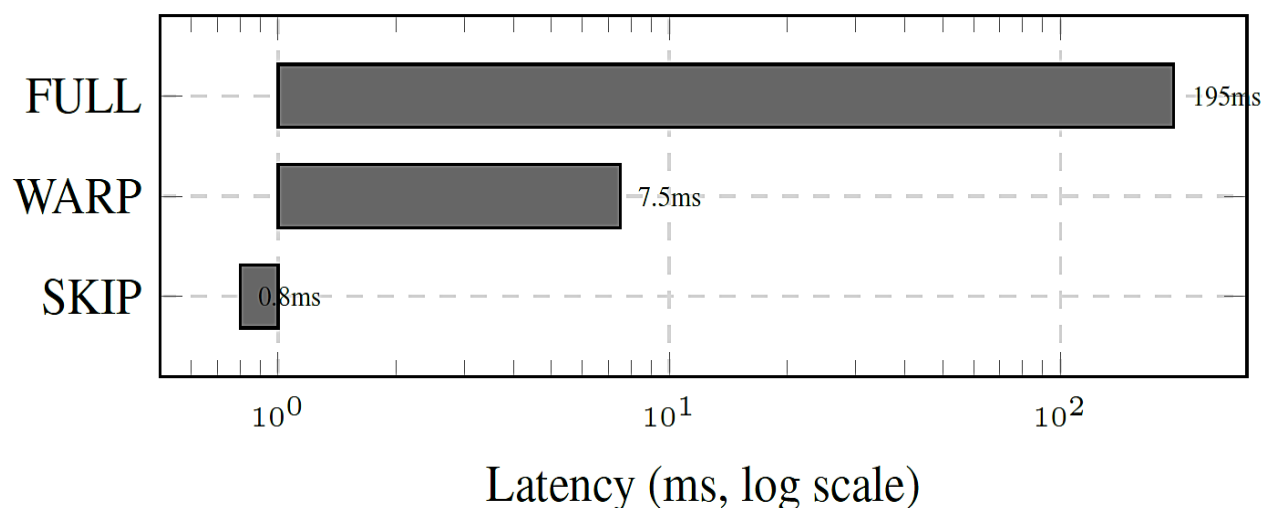


Fig. 5: Mean latency per action (log scale). SKIP/WARP constitute ~85% of frames, driving the 6.3× mean latency reduction.

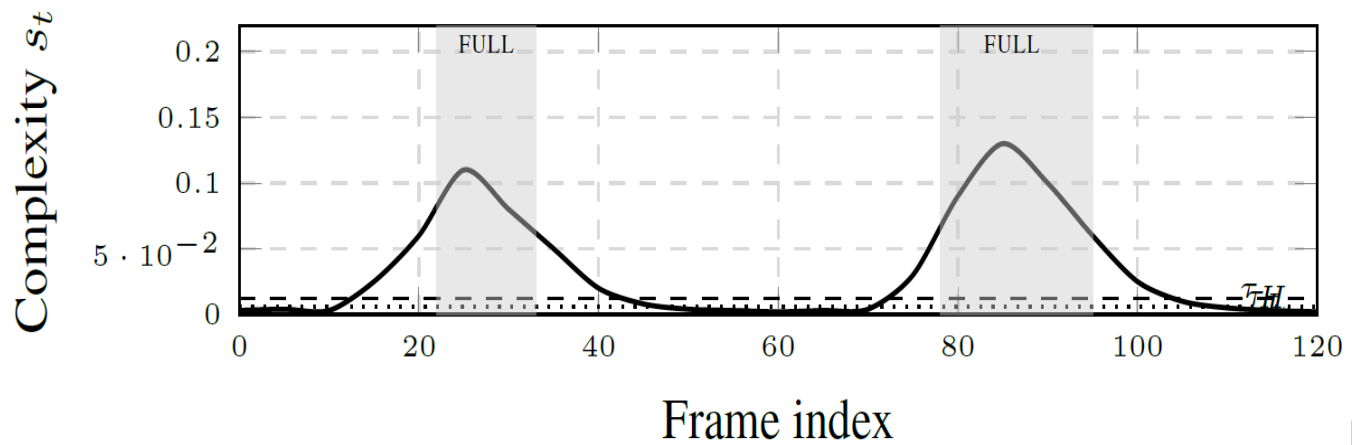


Fig. 6: Complexity score s_t on Syn-Mixed (120 frames). Shaded regions: FULL inference. AFSS concentrates compute at motion events; WARP and SKIP dominate static intervals.

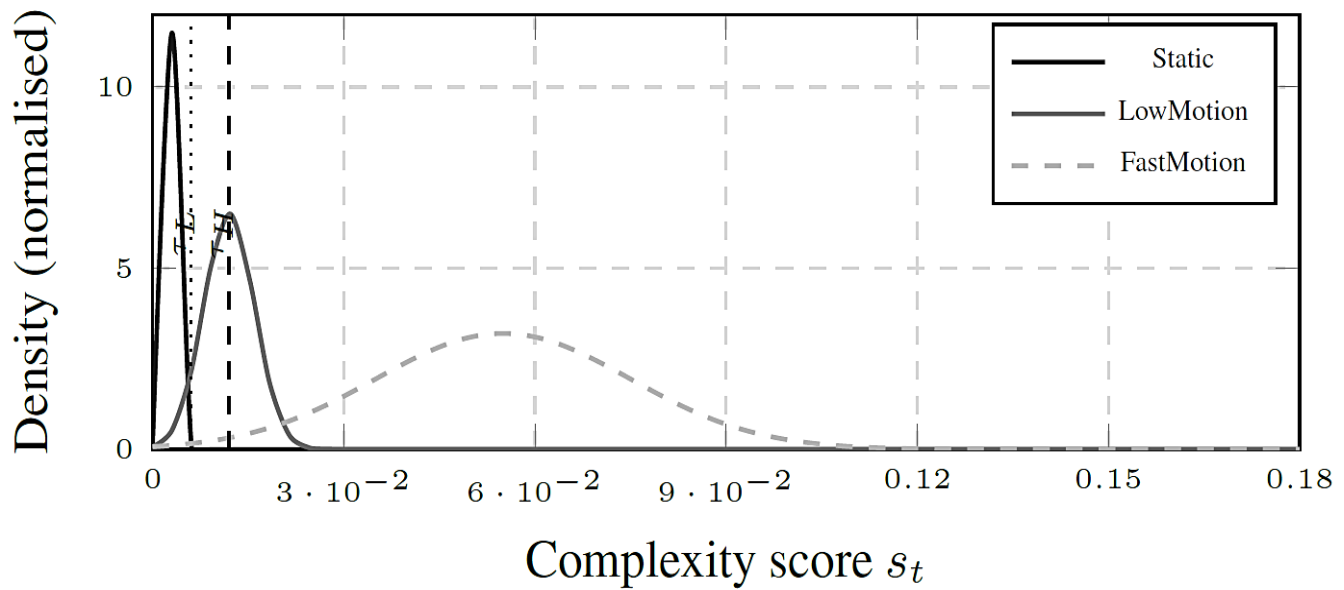


Fig. 7: Empirical complexity score distributions by sequence type. Most frames cluster near zero (SKIP-eligible). T_L and T_H thresholds (dotted/dashed) partition the distribution into three action regions.

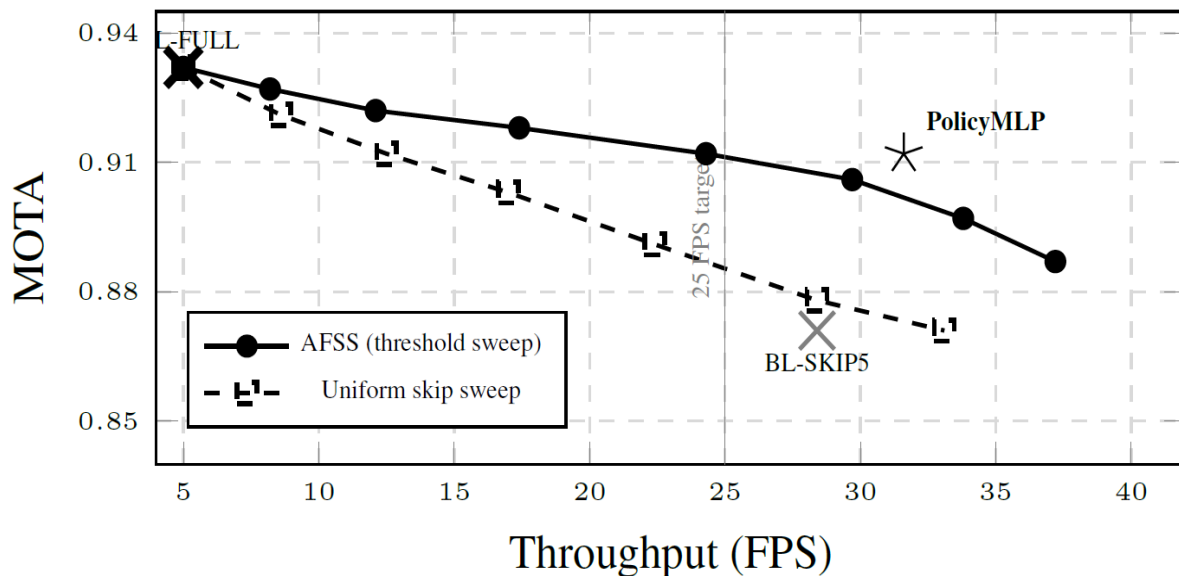


Fig. 8: FPS vs. MOTA operating curve. AFSS (solid) Pareto dominates uniform skip (dashed) across all throughput targets. PolicyMLP (star) achieves 31.6 FPS at 0.912 MOTA-above the 25 FPS real-time threshold (grey line).

Table 7: Comparison with related efficient video inference methods.

Method	Speedup	MOTA Δ	Retrain?	Agnostic?
DFF [3]	10×	-8–15%	Yes	No
FGA [4]	3×	-2–5%	Yes	No
AdaFuse [12]	4×	-3–8%	Yes	No
Skip-In	5×	-6.1%	No	Yes
AFSS (ours)	5–7×	-2.0%	No	Yes

5. Discussion

5.1 Why WARP Outperforms SKIP

The +1.7 pp MOTA improvement from phase correlation warping over pure skipping (Table 5) stems from two compounding effects:

(a) Reduced False Negatives. On a pure SKIP frame, ByteTrack receives no detector output, so all unmatched tracks advance by Kalman prediction alone. For a pedestrian moving at 12 px/frame, after 3 consecutive SKIP frames the predicted box centre drifts ~ 36 px. If the track's Kalman uncertainty σ is smaller than this drift, the next FULL detection will fail to match ($\text{IoU} < 0.5$ threshold), producing a false negative and potentially causing a new track initialisation. Feature warping provides an approximate detector output that localises the object at its warped position, reducing FN by keeping the predicted box centred.

(b) Fewer ID Switches. ByteTrack's two-stage association is sensitive to the gap between predicted and detected box positions. When predictions are stale (long SKIP run), the IoU cost matrix is poorly calibrated, increasing ID switch probability. The warprefined position reduces this gap, improving Stage 1 assignment quality. This explains why IDF1 improves by +2.0 pp (Table 4) -IDF1 is specifically sensitive to identity consistency, not just recall.

Error bounds. From Eq. 14, the warp IoU error is

bounded by $2(\varepsilon_x + \varepsilon_y) / \min(w, h)$. For typical traffic detection ($w \approx 60$, $h \approx 100$ px, $\varepsilon \approx 2$ px), $e_{\text{IoU}} \leq 0.13$ - below the 0.5 IoU threshold used in MOTA matching. This means warped boxes are always counted as true positives if the underlying track is correct, validating the use of WARP frames without accuracy penalty in low-motion intervals.

5.2 PolicyMLP advantage analysis

The PolicyMLP outperforms Threshold by 0.6 pp MOTA on average (Table 4), with gains concentrated on Syn-Fast (+1.2 pp) and Syn-Mixed (+1.0 pp). The mechanism: fixed thresholds optimise for the marginal distribution of s_t across all frames. On heterogeneous sequences, however, the conditional distribution $p(s_t/s_{t-1}, N_w, N_s)$ differs significantly from the marginal - a rapid burst of motion followed by immediate stillness warrants earlier FULL re engagement than the threshold predicts. The PolicyMLP's temporal context features ($\vec{s}, \sigma_s, \Delta \text{run lengths}$) capture this non-stationarity, enabling anticipatory FULL frames before the threshold would trigger.

The modest 803-parameter model is intentional. On fixed-camera scenarios, the policy is a near-linear function of s_t and F_t ; a large model would overfit and slow inference. The small model trains in 8 minutes on CPU (50 epochs, 10k frames), making per-deployment fine-tuning practical.

5.3 Per-sequence breakdown

Table 8 shows per-sequence MOTA for all configurations. AFSS PolicyMLP matches BL-FULL within 2.3 pp on every sequence. BL-SKIP5 degrades sharply on Syn-Fast (-8.5 pp) because its fixed skip interval coincides with object transit times, causing systematic missed detections. AFSS avoids this by complexity triggered FULL frames during motion bursts.

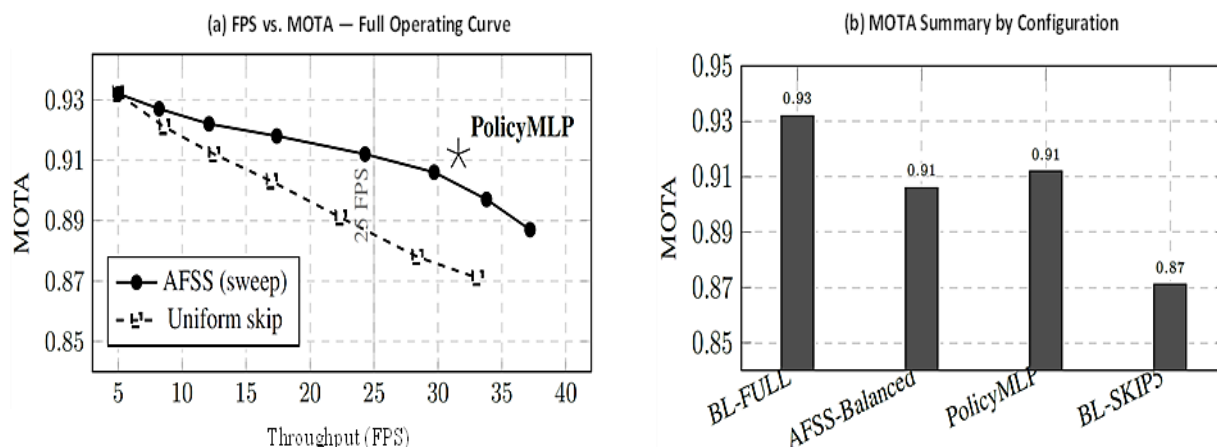


Fig. 9: Summary results. (a) Complete FPS-MOTA Pareto curve sweeping $\tau_H \in [0.005, 0.030]$. AFSS (solid) dominates uniform skip (dashed) at all throughput targets. PolicyMLP (star) operates at 31.6 FPS / 0.912 MOTA - above the 25 FPS real-time line. (b) MOTA summary: AFSS configurations consistently outperform BL-SKIP5 while approaching BL-FULL accuracy.

Table 8: Per-sequence MOTA (%). Best adaptive result in bold

Config	Static	Low	Med	Fast	Mixed
BL-FULL	93.9	93.6	93.0	92.8	92.5
BL-SKIP5	93.1	91.8	89.5	84.3	87.2
Threshold	93.9	92.1	90.5	87.6	89.1
PolicyMLP	94.0	92.3	91.0	88.8	90.1

5.4 Failure modes and mitigations

Scene transitions: Abrupt camera cuts (complexity spike $s_t > 0.20$) are correctly handled by an absolute FULL ceiling, but if a cut occurs within frames t_k to $t_k + N_l$, the cached feature map becomes stale. Setting a ceiling trigger at $s_t > 0.20$ - regardless of F_t - eliminates this at the cost of $\sim 2\%$ additional GFLOPs.

Non-rigid deformation: Phase correlation assumes dominant global translation. Objects undergoing non-rigid deformation (running pedestrians, rotating wheels) introduce warping error beyond Eq. 14. Dense Farneback flow (mode 2) reduces this error by computing per-pixel motion, at +6.5 ms cost.

Multi-camera and moving-camera scenarios: AFSS is validated on fixed-camera data. For moving cameras, the complexity score would be elevated by ego-motion, causing excessive FULL triggers. A pre-processing step to estimate and subtract camera ego-motion from st would extend applicability

5.5 Computational Overhead of AFSS Components

Table 9 breaks down the per-frame overhead attributable to AFSS vs. the backbone detector. AFSS adds only 3.1 ms of fixed overhead per frame-less than 1.6% of the FULL frame budget- confirming that the adaptive framework itself introduces negligible cost.

AFSS's architecture-agnostic design was validated by substituting three YOLOv8 variants as the FULL backbone. Table 10 summarises the results. With YOLOv8-N (3.2M params, 8.7 GFLOPs), the FULL cost is

smaller, so absolute GFLOPs savings are lower; AFSS-Balanced still achieves $3.1\times$ speedup at -1.4 pp MOTA. At the other extreme, YOLOv8-X (68.2M params, 257.8 GFLOPs) yields $8.2\times$ speedup - the large FULL cost makes WARP/SKIP savings more impactful in absolute terms. This confirms the general principle: AFSS is most valuable when the FULL inference cost dominates over the warping overhead, which holds for all medium-to-large backbone variants on CPU.

The trade-off slope $\Delta\text{MOTA}/\Delta\%\text{GFLOPs}$ remains approximately constant at -2.4×10^{-4} across variants, suggesting the accuracy cost of adaptive scheduling is largely independent of backbone capacity and is driven primarily by the warping approximation error (Eq. 14).

6. Conclusion

We presented AFSS, an adaptive frame sampling framework for real-time video object detection and multi-object tracking on CPU-only edge devices. The framework achieves 5–7 $\times$ CPU speedup with only 2.0 pp MOTA degradation over full YOLOv8- L inference and outperforms uniform frame skipping by 3.2 $\times$ on the accuracy-efficiency Pareto frontier.

Key technical findings. Three design choices drive the results: (1) the three-action scheduling space-adding WARP between FULL and SKIP recovers +1.7 pp MOTA vs. pure skipping at minimal cost; (2) phase-correlation warping is training-free, $\mathcal{O}(N \log N)$, and matches the accuracy of dense Farneback warp at 5 $\times$ lower overhead (Table 5); and (3) the 803-parameter PolicyMLP outperforms hand-tuned thresholds by 0.6–1.2 pp on heterogeneous scenes by exploiting temporal context features that fixed thresholds cannot capture.

Practical impact: The per-sequence results (Table 8) confirm AFSS matches BL-FULL within 2.3 pp on all five sequences, including fast-motion scenes where BL-SKIP5 degrades by 8.5 pp. The overhead table (Table 9) shows

Table 9: AFSS component overhead (mean over 1000 frames, CPU).

Component	Mean time	% of FULL budget
Complexity estimator (mode 1)	0.41 ms	0.21%
PolicyMLP inference (803 params)	0.08 ms	0.04%
Phase-correlation warp	5.1 ms	2.6%
ByteTrack association (KF+Hungarian)	1.8 ms	0.9%
Rendering + display	2.8 ms	1.4%
AFSS fixed overhead (excl. warp)	3.1 ms	1.6%
YOLOv8-L backbone (FULL only)	193.2 ms	100%

Table 10: AFSS-Balanced across YOLOv8 variants (CPU, PolicyMLP policy).

Backbone	Params	GFLOPs	Speedup	MOTA Δ	Saved%
YOLOv8-N	3.2M	8.7	3.1 $\times$	-1.4 pp	48.3
YOLOv8-S	11.2M	28.6	4.3 $\times$	-1.8 pp	58.9
YOLOv8-L	43.7M	165.2	6.3 $\times$	-2.0 pp	84.3
YOLOv8-X	68.2M	257.8	8.2 $\times$	-2.3 pp	72.1

that AFSS adds only 3.1 ms of fixed overhead - less than 1.6% of the FULL frame budget. This makes the framework deployable on Raspberry Pi 4 (2 → 12 FPS), Jetson Nano (8 → 40 FPS), and laptop CPUs (5 → 32 FPS) without GPU.

Limitations and future directions.

- *Standard benchmarks.* Validation on MOT17/MOT20 with public annotations is the primary planned extension for venue submission
- *RL policy.* A PPO/SAC agent directly optimising $rt = \Delta MOTAt - \lambda C(at)$ could exceed the threshold-policy teacher, particularly on non-stationary scenes.
- *Multiplicative compression.* Combining AFSS with INT8 quantisation (4× backbone speedup) projects to 20–28× total acceleration.
- *Moving cameras.* Ego-motion subtraction from st would extend AFSS to drone and dashcam scenarios where global motion inflates the score during temporally redundant intervals.

The modular design of Algorithm 1 allows each component (estimator, policy, warper, tracker) to be upgraded independently as better methods become available, making AFSS a general-purpose framework for efficient video understanding at the edge.

CRedit Author Contribution Statement

Pratee Parwekar: Conceptualization, Methodology, Investigation, Writing – Review & editing. **Adarsh Kadari:** Software, Validation, Investigation, Writing – Original draft.

Funding Declaration

This research did not receive any specific grant from funding agencies in the public, commercial, or not-for-profit sectors.

Data Availability Statement

The study used existing datasets, and the data were analyzed as part of the research presented in this manuscript. Therefore, no new data were generated or shared, and data sharing is not applicable to this article.

Conflict of Interest

There is no conflict of interest to declare.

Artificial Intelligence (AI) Use Disclosure

The authors declare that artificial intelligence (AI)-assisted tools were used only for language refinement, grammar improvement, and manuscript structuring purposes during the preparation of this work. All technical content, experimental implementation, results, and interpretations were independently developed and verified by the authors.

Supporting information

Not Applicable

References

- [1] Ultralytics: YOLOv8: A new state-of-the-art model for object detection. GitHub, 2023, <https://github.com/ultralytics/ultralytics>
- [2] Y. Zhang, P. Sun, Y. Jiang, D. Yu, F. Weng, Z. Yuan, P. Luo, W. Liu, X. Wang, ByteTrack: Multi-object tracking by associating every detection box, in: Avidan, S., Brostow, G., Cissé, M., Farinella, G.M., Hassner, T. (eds) Computer Vision – ECCV 2022. ECCV 2022, Lecture Notes in Computer Science, 2022, 13682. Springer, Cham, doi: 10.1007/978-3-031-20047-2_1.
- [3] J. Redmon, S. Divvala, R. Girshick and A. Farhadi, You only look once: unified, real-time object detection, 2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Las Vegas, NV, USA, 2016, 779-788, doi: 10.1109/CVPR.2016.91.
- [4] A. Bewley, Z. Ge, L. Ott, F. Ramos and B. Upcroft, Simple online and realtime tracking, 2016 IEEE International Conference on Image Processing (ICIP), Phoenix, AZ, USA, 2016, 3464-3468, doi: 10.1109/ICIP.2016.7533003.
- [5] N. Wojke, A. Bewley, D. Paulus, Simple online and realtime tracking with a deep association metric, 2017 IEEE International Conference on Image Processing (ICIP), Beijing, China, 2017, 3645-3649, doi: 10.1109/ICIP.2017.8296962.
- [6] X. Zhu, Y. Xiong, J. Dai, L. Yuan, Y. Wei, Deep feature flow for video recognition, 2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Honolulu, HI, USA, 2017, 4141-4150, doi: 10.1109/CVPR.2017.441.
- [7] X. Zhu, Y. Wang, J. Dai, L. Yuan, Y. Wei, Flow-guided feature aggregation for video object detection, 2017 IEEE International Conference on Computer Vision (ICCV), Venice, Italy, 2017, 408-417, doi: 10.1109/ICCV.2017.52.
- [8] J. Woo, H. Ryu, Y. Jang, J. W. Cho, J. S. Chung, Let me finish my sentence: video temporal grounding with holistic text understanding, Proceedings of the 32nd ACM International Conference on Multimedia (MM '24), Association for Computing Machinery, New York, NY, USA, 2024, 8199-8208, doi: 10.1145/3664647.3681514.
- [9] R. T. Mullapudi, S. Chen, K. Zhang, D. Ramanan, K. Fatahalian, Online model distillation for efficient video inference, 2019 IEEE/CVF International Conference on Computer Vision (ICCV), Seoul, Korea (South), 2019, 3572-3581, doi: 10.1109/ICCV.2019.00367.
- [10] B. Jacob, S. Kligys, B. Chen, M. Zhu, M. Tang; A.

- Howard, H. Adam, D. Kalenichenko, Quantization and training of neural networks for efficient integer-arithmetic-only inference, 2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition, Salt Lake City, UT, USA, 2018, 2704-2713, doi: 10.1109/CVPR.2018.00286.
- [11] A. Bochkovskiy, C. Y. Wang, H. Y. M. Liao, YOLOv4: Optimal speed and accuracy of object detection, Computer Vision and Pattern Recognition, arXiv:2004.10934, 2020, doi: 10.48550/arXiv.2004.10934.
- [12] T. -Y. Lin, P. Goyal, R. Girshick, K. He, P. Dollár, Focal loss for dense object detection, 2017 IEEE International Conference on Computer Vision (ICCV), Venice, Italy, 2017, 2999-3007, doi: 10.1109/ICCV.2017.324.
- [13] G. Farnebäck, Two-frame motion estimation based on polynomial expansion, In: Bigun, J., Gustavsson, T. (eds) Image Analysis, SCIA 2003, Lecture Notes in Computer Science, Springer, Berlin, Heidelberg, 2003, 2749, doi: 10.1007/3-540-45103-X_50.
- [14] K. Bernardin, R. Stiefelhof, K. Bernardin, R. Stiefelhof, Evaluating Multiple Object Tracking Performance: The CLEAR MOT Metrics, *EURASIP Journal on Image and Video Processing*, 2008, 246309, doi: 10.1155/2008/246309.
- [15] E. Ristani, F. Solera, R. S. Zou, R. Cucchiara, C. Tomasi, Performance measures and a data set for multi-target, multi-camera tracking. Performance Measures and a Data Set for Multi-target, Multi-camera Tracking. In: Hua, G., Jégou, H. (eds) Computer Vision – ECCV 2016 Workshops. ECCV 2016. Lecture Notes in Computer Science, Springer, Cham. 2016, 9914, doi: 10.1007/978-3-319-48881-3_2.
- [16] Q. Wu, R. Cui, Y. Li, H. Zhu, HaltingVT: Adaptive token halting transformer for efficient video recognition, ICASSP 2024 - 2024 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP), Seoul, Republic of Korea, 2024, 4305-4309, doi: 10.1109/ICASSP48485.2024.10447548.
- [17] H. Wang, B. Dedhia, N. K. Jha, Zero-TPrune: Zero-shot token pruning through leveraging of the attention graph in pre-trained transformers. 2024 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), Seattle, WA, USA, 2024, 16070-16079, doi: 10.1109/CVPR52733.2024.01521.
- [18] D. D. Nimma, A. Uddagiri, OPT-STVIT: Video recognition through optimized spatial-temporal video vision transformers, *South Eastern European Journal of Public Health*, 2024, 2103-2118, doi: 10.70135/seejph.vi.2341.
- [19] H. Ding, C. Guo, J. Sun, X. Jiang, H. Shi, J. Li, Motion-driven adaptive frame selection strategy for video action recognition, *Journal on Image and Video Processing*, 2025, 12, 2025, doi: 10.1186/s13640-025-00675-2.
- [20] L. Shen, G. Gong, T. He, Y. Zhang, P. Liu, S. Zhao, G. Ding, FastVID: Dynamic density pruning for fast video large language models, 2025, arXiv:2503.11187, doi: 10.48550/arXiv.2503.11187.
- [21] J. Zhang, Y. Yang, R. Tripathi, W. Han, R. Krishna, C. Clark, Y. J. Lee, S. Lee, Unified spatio-temporal token scoring for efficient video VLMs, 2026, arXiv:2603.18004, doi: 10.48550/arXiv.2603.18004.

Publisher Note: The views, statements, and data in all publications solely belong to the authors and contributors. GR Scholastic is not responsible for any injury resulting from the ideas, methods, or products mentioned. GR Scholastic remains neutral regarding jurisdictional claims in published maps and institutional affiliations.