

DARC-OR: Dual-Signal Adaptive Recalibration and Concept-Drift Response for Capacity-Constrained Operational-Risk Early Warning

Ethan Mercer

Manning College of Information and Computer Sciences, University of Massachusetts Amherst, MA, the United States

Abstract: Operational-risk early-warning models in financial institutions operate under nonstationarity, extreme class imbalance, delayed outcome confirmation, and finite investigation capacity. This paper proposes DARC-OR, a drift-response framework that combines an immediately available Jensen-Shannon divergence signal on model scores with a delayed supervised EWMA signal on balanced log loss. The controller distinguishes short-lived distribution change from persistent or performance-relevant change, using recalibration for the former and recency-weighted adaptive-window retraining for the latter. A controlled 120-day transaction stream benchmark was executed for five independent seeds, totaling 720,000 transactions and 6,839 positive events, with abrupt, gradual, and recurring concept changes, a three-day label delay, and a 0.5% daily review capacity. Across the five runs, DARC-OR achieved average precision 0.1719 ± 0.0280 , exceeding a delayed-loss trigger by 21.55% and a static model by 38.22%; recall at capacity was 0.1423 ± 0.0175 and Brier score was 0.008374 ± 0.000305 . However, its normalized exposure scenario cost (0.7221) was worse than both delayed-loss (0.6888) and static (0.6540) policies, and recurrence exposed substantial forgetting. The results therefore support dual-signal adaptation as a ranking-and-calibration improvement under drift, while showing that drift response, action cost, and recurring-concept memory must remain separately governed.

Keywords: Concept drift, operational risk, early warning, delayed labels, calibration, cost-sensitive learning, streaming machine learning, financial institutions.

1. Introduction

Operational risk is not a static classification problem. Financial institutions must identify changing process, people, systems, and external-event risks while maintaining governance, validation, and human escalation. The Basel Committee explicitly links operational-risk management to governance, monitoring, control, change management, and information and communication technology [1]. Model-risk expectations likewise emphasize a risk-based lifecycle in which model development, validation, monitoring, and change control remain distinct responsibilities [2]. For AI-enabled controls, the NIST AI Risk Management Framework adds validity, reliability, transparency, explainability, and accountability as connected trustworthiness concerns [3]. These requirements make temporal behavior part of the risk model itself rather than a secondary reporting issue.

A strong recent contribution in this direction is XOR-EW, which integrates chronological controls, calibrated probabilities, capacity-aware action, uncertainty routing, and auditable governance for transaction-level operational-risk early warning [4]. DARC-OR adopts the same constructive premise that a usable warning system is a governed decision pipeline, and extends that direction by making the temporal control layer explicitly responsive to concept drift under delayed labels. The purpose is not to replace calibration, policy optimization, or human review, but to determine when the predictive component should be recalibrated or relearned as the stream changes.

Concept drift occurs when the joint distribution relevant to prediction changes over time. Surveys distinguish abrupt, gradual, incremental, and recurring forms, and emphasize that detection and adaptation are separate design problems [5-7]. In financial streams the difficulty is amplified because labels

are rarely instantaneous: an incident can be confirmed days or weeks after the score was produced. A detector based only on labels can therefore react too late, while a detector based only on covariates or scores can overreact to harmless distribution change. The practical question is consequently not just whether drift exists, but what evidence is available now, how persistent it is, and which intervention is proportionate.

This paper studies a focused, reproducible version of that question. DARC-OR combines two evidence channels: (1) an unlabeled Jensen-Shannon divergence (JSD) between a reference score distribution and the current-day score distribution; and (2) a delayed supervised exponentially weighted moving average (EWMA) of class-balanced log loss. A dual-signal controller maps evidence to either recalibration or adaptive-memory retraining. The model then feeds a capacity-constrained, exposure-sensitive triage rule. The architectural novelty is the response logic and its governance separation; JSD, EWMA, Platt calibration, and gradient boosting are established components.

The empirical contribution is a controlled stream benchmark designed to expose failure modes that a two-day static dataset cannot identify. Five independent 120-day streams contain an abrupt A-to-B change, a gradual B-to-C transition, a recurrence of A, and an additional covariate shift that is not itself a label-concept change. All labels are delayed by three days. The executed benchmark contains 720,000 transactions and 6,839 positive events, and all numbers reported below are computed from the supplied run artifacts. The study makes four contributions: (i) dual-signal drift evidence under delayed feedback; (ii) tiered recalibration versus adaptive retraining; (iii) capacity-aware evaluation that keeps predictive adaptation separate from policy cost; and (iv) a fully logged prequential protocol that reports both gains and failure modes, including recurring-concept

forgetting.

2. Related Work and Research Gap

2.1. Concept Drift Detection and Adaptation

Concept-drift research has developed a rich taxonomy of detection and response mechanisms. Gama et al. organize the field around drift types, detection, and adaptation [5], while Lu et al. formalize detection, understanding, and adaptation as related but distinct activities [6]. Webb et al. further separate changes in class priors, covariates, and posterior relationships, which is important because not every observable shift implies a deterioration of the target relationship [7]. Adaptive Windowing (ADWIN) provides a statistically motivated variable-window change detector [8]. Sethi and Kantardzic show why unlabeled drift monitoring can be useful but vulnerable to false alarms when distribution change is not performance relevant [9]. Under simultaneous class imbalance and concept drift, the evaluation problem becomes harder because minority-class evidence is sparse and delayed [10].

DARC-OR is positioned between purely supervised and purely unlabeled detectors. It does not claim that JSD is a formal test for posterior drift, nor that an EWMA ratio isolates the source of degradation. Instead, each signal is assigned an operational role: JSD supplies immediate evidence that the score-generating environment moved, and delayed balanced loss supplies evidence that predictive utility has degraded after labels arrive. Persistence and cooldown rules then determine response severity. This hybrid evidence strategy is deliberately simple enough to reconstruct from an audit log.

2.2. Streaming Fraud Detection as an Operational-Risk Proxy

Transaction fraud is only one observable proxy for operational loss and does not represent the full Basel operational-risk taxonomy. It is nevertheless a useful testbed because it combines rarity, asymmetric consequences, adaptive adversaries, high arrival rates, and delayed verification. Dal Pozzolo et al. explicitly studied concept-drift adaptation with delayed supervised information in card fraud [11], and later developed a realistic learning strategy that incorporates class imbalance, concept drift, and verification latency [12]. SCARFF demonstrated scalable streaming fraud analytics [13], while hybrid supervised-unsupervised approaches exploit anomaly evidence alongside learned fraud patterns [14]. The Fraud Detection Handbook provides reproducible designs for realistic evaluation and simulation [15].

Public financial datasets, however, rarely provide known timestamps for abrupt, gradual, harmless, and recurring drift. Benchmarking studies warn that real-world streams often have uncertain change points and hidden preprocessing histories, which makes detector-delay comparisons difficult [16]. For this reason, the present paper uses a controlled synthetic stream rather than inventing drift labels for a public static dataset. The synthetic design is not presented as institutional realism; it is a stress test in which the ground-truth change schedule is known and every generated observation is reproducible from code and seed.

2.3. Imbalance, Calibration, and Decision Cost

Operational early warning is a rare-event problem. General imbalanced-learning research shows that accuracy can be

misleading when the positive class is small [17]. Precision-recall analysis is therefore emphasized for rare-event ranking [18], while the mathematical relation between ROC and precision-recall spaces explains why prevalence changes the apparent usefulness of a classifier [19]. DARC-OR uses average precision (AP) as the principal ranking metric and also reports AUROC, but does not infer an action policy from either threshold-free metric alone.

Probability quality is equally important because a drift response may preserve ranking while changing score scale. Brier introduced the squared probability score used here for calibration assessment [20]. Later work demonstrated that high predictive accuracy and good probability estimates are separate properties [21], [22]. Platt-style logistic calibration remains a practical post-hoc transformation when a dedicated calibration sample is available [23]. DARC-OR therefore permits calibration-only responses rather than treating every score-distribution shift as evidence that the base learner must be replaced.

Finally, operational decisions are cost-sensitive. Elkan formalized the foundations of cost-sensitive learning [24], and example-dependent work in financial fraud shows why the consequence of a missed event can vary with transaction exposure [25], [26]. The present benchmark applies a transparent exposure scenario and a hard daily investigation capacity, but it treats this objective as a policy evaluation layer rather than a monetary loss estimate. This separation becomes central to interpreting the results: better AP does not automatically imply lower operational cost.

3. Problem Formulation

Let transactions arrive in daily batches $B_t = \{(x_i, a_i)\}$ with feature vector x_i and transaction amount a_i . The binary outcome y_i in $\{0,1\}$ becomes available only after a fixed verification delay L . The predictor at day t produces a raw score $s_t(x)$ and a calibrated probability $p_t(x)$. Concept drift means that the data-generating process $P_t(X,Y)$ is not stationary. We distinguish observable score-distribution change from performance-relevant degradation because $P_t(X)$ can move without necessarily changing the decision boundary.

The operational policy has a finite daily review fraction κ . In each batch, DARC-OR forms an exposure-prioritized score $q_i = p_t(x_i)[1 + a_i / \text{median}(a \text{ in } B_t)]$ and sends the top $\text{ceil}(\kappa|B_t|)$ transactions to review. This is a ranking policy under a hard workload cap, not an estimate of accounting loss. In the main experiment $\kappa = 0.005$, so each day at most 0.5% of transactions is selected.

For retrospective evaluation, a false alert costs one unit and a missed positive costs $1 + a_i / \text{median}(a)$ over the evaluation stream. Normalized cost divides the realized scenario cost by the cost of clearing every transaction. Values below one indicate an improvement over the all-clear scenario, but comparisons between models remain sensitive to the chosen coefficients. Consequently, the optimization target of the drift controller is not defined to be this test cost. This prevents a post hoc claim that a detector is superior merely because one policy coefficient favors it.

4. Darc-Or Framework

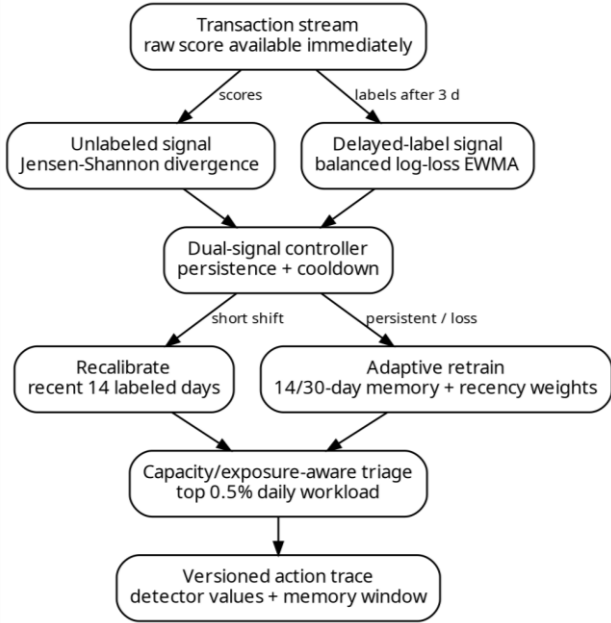


Figure 1. DARC-OR dual-signal drift response and governance flow

4.1. Initialization and Base Learner

The first 20 days form the initial training window, days 20-24 form a separate calibration window, and days 25-29 establish detector baselines. The base learner is LightGBM [27] with 100 trees, learning rate 0.045, 23 leaves, minimum child size 60, 0.85 row and column subsampling, L2 regularization 1.0, and class_weight="balanced". These values are fixed for all strategies. Raw scores are mapped through a logistic calibrator fitted on the dedicated calibration window. The evaluation begins on day 30, after all initial model, calibration, and detector settings are fixed.

4.2. Immediate Unlabeled Drift Signal

For each day, the current raw-score histogram is compared with a reference distribution built from the calibration period or the most recent labeled reference window. DARC-OR uses Jensen-Shannon divergence, a symmetric and finite information divergence [28]. With $M = (P+Q)/2$,

$$\text{JSD}(P, Q) = 0.5 \text{KL}(P \parallel M) + 0.5 \text{KL}(Q \parallel M) \quad (1)$$

Where the histograms use decile boundaries estimated from the reference scores and a small 0.5 pseudocount. The baseline threshold is $\max(0.004, \text{median}(\text{JSD}_{25:29}) + 5 \times 1.4826 \times \text{MAD})$. This is a robust engineering threshold, not a calibrated hypothesis-test level. An exceedance increments an unlabeled persistence counter; a non-exceedance resets it.

4.3. Delayed Supervised Performance Signal

When labels for day $t-L$ arrive, DARC-OR computes a balanced log loss so that the rare positive class and the negative class contribute equal class-level weight. If L_+ and L_- are the mean log losses in the two classes, the daily loss is $\text{ell}_t = 0.5 L_+ + 0.5 L_-$. The loss is then smoothed using the classical EWMA control idea [29]:

$$z_t = 0.35 \text{ell}_t + 0.65 z_{t-1} \quad (2)$$

A supervised warning is active when z_t divided by its robust recent baseline exceeds 1.35 for three consecutive labeled updates. The baseline is refreshed after a successful

model action using the median daily balanced loss over recent labeled data. This signal therefore arrives later than JSD but is closer to the actual predictive objective.

4.4. Tiered Response and Adaptive Memory

The controller uses severity and persistence rather than a single retraining trigger. Two consecutive JSD exceedances can produce a calibration-only response if at least ten days have elapsed since the previous DARC-OR action. Five consecutive JSD exceedances escalate to retraining. A persistent supervised-loss warning also escalates to retraining. The ten-day action separation is intentionally conservative relative to the three-day label delay and reduces repeated interventions on adjacent batches. Calibration-only updates use the most recent 14 labeled days.

For retraining, DARC-OR chooses between 14-day and 30-day memory. Each candidate window is split chronologically, trained on the earlier 75%, and evaluated on the later 25% using an operational validation utility $U(w) = 0.45 \text{AP} + 0.35 \text{R}_{\text{cap}} + 0.20(1 - \text{C}_{\text{norm}})$. The chosen learner is then fitted with exponential recency weights $\exp(-\text{age}/\tau)$, where $\tau = \max(5, w/3)$, and calibrated on the most recent 20% of the chosen window. This adaptive memory is designed to trade fast forgetting against sample size without consulting future labels.

4.5. Action Trace and Governance Boundary

Every intervention is recorded with stream day, action type, JSD value, delayed loss ratio, and selected memory window. The trace makes four objects independently reviewable: detector state, model state, calibrator state, and decision policy. DARC-OR intentionally does not auto-tune the 0.5% review cap or the exposure scenario after seeing test outcomes. A financial institution could therefore challenge detector thresholds or retraining frequency without silently changing alert economics.

5. Experimental Design

5.1. Controlled Stream Generator

The benchmark generates 120 days with 1,200 transactions per day. Each seed initializes 3,000 latent customer profiles and 500 latent terminal profiles. The model observes 11 features: log amount, amount-to-customer-baseline ratio, sine and cosine hour encodings, weekend indicator, cross-border indicator, device novelty, log geographic distance, log transaction velocity, terminal risk, and cash-like indicator. Labels are sampled from a logistic event mechanism whose coefficients change by regime. The generator also inserts a temporary covariate shift around days 72-92 by changing legitimate cross-border, device, distance, and cash-like behavior independently of the event mechanism. This creates conditions in which an unlabeled detector can respond to a shift that is not by itself evidence of target degradation.

Table I. Controlled Temporal Benchmark

Stage	Days	Role / regime
Train	0-19	Initial regime A; model fit
Calibrate	20-24	Platt calibration + score reference
Baseline	25-29	JSD/loss detector baselines
Evaluate	30-44	Regime A
Abrupt	45-69	A -> B at day 45
Gradual	70-84	B -> C transition
C	85-94	Regime C
Recurrence	95-119	Return to regime A

A separate covariate shift is strongest near day 82; labels are delayed by three days.

Five executed seeds (2023-2027) produce 144,000 observations each, for 720,000 total transactions. Positive counts are 1,392, 1,361, 1,305, 1,392, and 1,389, respectively; the pooled event rate is 0.9499%. These values are outputs of the supplied generator, not imposed class quotas. The event prevalence is deliberately higher than the ULB two-day fraud dataset because each post-drift regime must contain enough delayed positives to evaluate a detector within a 120-day controlled experiment.

5.2. Baselines and Prequential Protocol

All strategies share the same initial LightGBM model, Platt calibrator, features, label delay, and daily capacity. Static performs no update. Periodic30 retrains a fixed 30-day window every 14 days. Loss-trigger retrains a fixed 30-day window when the delayed supervised warning fires. JSD-trigger retrains a fixed 30-day window on the unlabeled JSD signal. DARC-OR is the only strategy that distinguishes recalibration from retraining and selects 14- versus 30-day adaptive memory with recency weighting.

Evaluation is prequential from day 30 through day 119. Predictions for the current day are produced before that day can enter any labeled training set. At day t , retraining can use only records with day $\leq t-3$. Final labels are used for retrospective metrics only after the online decision sequence has been fixed. The implementation records row-level predictions for the representative 2023 seed and daily, event, and summary files for all five seeds. No headline result is hard-coded in the analysis script.

5.3. Metrics

Ranking is measured by AP and AUROC; action quality is measured by precision and recall at the 0.5% daily capacity; calibration is measured by Brier score. We also report normalized exposure scenario cost, retraining and recalibration counts, unmatched retraining actions, and event-window delays. An action is matched to a known drift point when a retrain occurs within 14 days after the point. Because this is a matching convention rather than causal attribution, a zero-day match at the start of a gradual transition is not interpreted as proof of instantaneous statistical detection. Means and sample standard deviations are reported across the five executed seeds; with only five runs, the study avoids significance claims.

6. Results

6.1. Aggregate Predictive and Decision Performance

Table II. Five-Seed Mean Performance

Method	AP	AUC	P@cap	R@cap	Brier	Cost
Static	0.1243	0.8256	0.2530	0.1382	0.009528	0.6540
Periodic	0.1244	0.8065	0.2070	0.1131	0.008954	0.7759
JSD	0.1001	0.7838	0.2122	0.1159	0.009188	0.7713
Loss	0.1414	0.8511	0.2356	0.1287	0.008462	0.6888
DARC-OR	0.1719	0.8352	0.2607	0.1423	0.008374	0.7221

P@cap and R@cap use a 0.5% daily review capacity. Cost is the normalized exposure scenario; lower is better. DARC-OR standard deviations: AP 0.0280, AUC 0.0235, P@cap 0.0352, R@cap 0.0175, Brier 0.000305, cost 0.0311.

DARC-OR produced the highest mean AP, 0.1719, compared with 0.1414 for delayed-loss retraining and 0.1243 for the static model. The corresponding relative gains are 21.55% and 38.22%, respectively. DARC-OR AP exceeded delayed-loss AP in four of five seeds. Recall at capacity improved by 10.58% relative to delayed-loss and by 3.02% relative to static, while the Brier score was 12.11% lower than static. The best mean AUROC, however, belonged to delayed-loss retraining (0.8511), illustrating that DARC-OR is not uniformly best under every ranking summary.

The cost result is more important than a simple winner label. DARC-OR normalized cost was 0.7221, 4.84% higher than delayed-loss and 10.41% higher than static. In fact, static cost was lower than DARC-OR for every seed. The reason is structural: the benchmark policy prioritizes probability multiplied by exposure, whereas the adaptation utility balances AP, capture, and cost on recent validation data. A model can improve rare-event ranking while moving errors toward higher-exposure cases. The data therefore reject the claim that predictive drift adaptation automatically improves the chosen operating-cost scenario.

6.2. Temporal Behavior Under Drift

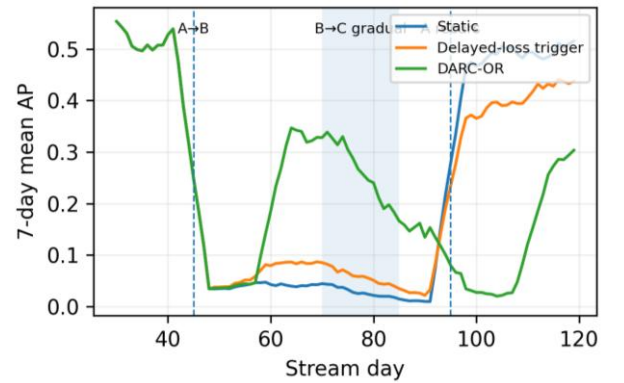


Figure 2. Seven-day mean AP across five seeds. Dashed lines mark the abrupt change at day 45 and recurrence at day 95; shading marks the gradual B-to-C transition.

Figure 2 shows the main temporal pattern. All methods perform similarly during regime A because they share the initial model. After the day-45 A-to-B change, the static model collapses and delayed-loss adaptation recovers only modestly. DARC-OR recovers substantially during B and the B-to-C transition because repeated evidence triggers short-memory retraining. During the C regime, all methods remain difficult because the positive rate is lower and the active

mechanism is substantially different from A. When A recurs at day 95, the static model immediately benefits from its retained A representation, whereas DARC-OR initially performs poorly because its recent-memory strategy has deliberately overwritten older A information.

Table III. Regime-Specific Mean Ap / Recall At Capacity

Regime	Static AP/R	Loss AP/R	DARC AP/R
A	0.4897 / 0.3149	0.4897 / 0.3149	0.4892 / 0.3149
B	0.0283 / 0.0053	0.0438 / 0.0042	0.1271 / 0.0541
B → C	0.0163 / 0.0040	0.0380 / 0.0040	0.2032 / 0.1124
C	0.0060 / 0.0000	0.0134 / 0.0040	0.0858 / 0.0900
A recur	0.4612 / 0.2764	0.3699 / 0.2462	0.1128 / 0.1701

Regime summaries pool each named evaluation segment within a seed and then average over five seeds.

Table III quantifies the trade-off. In regime B, DARC-OR AP is more than four times the static AP; in the gradual B-to-C segment, it is more than twelve times static; and in C, it maintains nonzero mean capture where static mean recall is zero. Yet recurring A reverses the ranking: static AP is 0.4612 while DARC-OR is 0.1128. This is a controlled example of catastrophic or recurring-concept forgetting. A production system facing seasonal or recurring incident patterns should therefore add concept memory, champion libraries, or ensemble reuse rather than relying only on short recent windows; stream-ensemble literature provides several directions for such extensions [30-34].

6.3. Drift Actions, Delay, and False-Alarm Burden

Table IV. Mean Action Burden and Event-Window Delay (Days)

Method	Retr.	Recal.	False	d45	d70	d95
Periodic	7.0	0.0	4.0	13	2	5
JSD	6.8	0.0	2.6	9*	4.6	2.4
Loss	2.6	0.0	0.2	5	-	5*
DARC-OR	6.4	0.6	2.4	5	0	5

False = retraining actions unmatched to a known drift window. *Delay mean uses only seeds with a matched event; JSD matched day 45 in one seed, and Loss matched recurrence in a subset. A dash denotes no matched day-70 retrain.

DARC-OR retrained 6.4 +/- 0.9 times and recalibrated 0.6 +/- 0.9 times per 90-day evaluation stream. It matched the day-45 abrupt event with a five-day delay in every seed and matched recurrence five days after day 95. The day-70 value is zero under the predefined event-window matcher because every seed has a retraining action on day 70. That result should be interpreted carefully: the controller already carried elevated loss state from regime B, so the day-70 action cannot be causally attributed only to the first instant of B-to-C drift. This is exactly why the supplied event log is retained rather than reporting a detector delay without context.

The action count also exposes a weakness. DARC-OR has 2.4 unmatched retraining actions on average, substantially above the delayed-loss trigger. The dual controller is therefore more responsive but also more intervention-heavy. In a real bank, retraining itself has validation, deployment, and

change-management costs. A reasonable deployment would place hard limits on retraining frequency, allow recalibration as a lower-impact action, and require challenger review when signal disagreement persists rather than automatically promoting each retrained model.

7. Discussion

7.1. What the Experiment Supports

The strongest supported claim is narrow: on this controlled delayed-label stream, combining immediate score-distribution evidence with delayed performance evidence produced more robust rare-event ranking through the new B and C regimes than static, periodic, unlabeled-only, or delayed-loss-only retraining. The advantage is not explained by a different base learner because every strategy uses the same LightGBM initialization. It arises from response timing, adaptive memory, recency weighting, and the ability to recalibrate without always retraining.

The second supported claim is that the two signals are complementary rather than interchangeable. JSD-only retraining reacted frequently but had the lowest mean AP, consistent with the possibility that distribution change is not always harmful. Delayed-loss retraining had the best mean AUROC and low action burden, but was weaker than DARC-OR in AP and recall during the new B/C regimes. The dual controller uses unlabeled persistence to obtain early evidence while requiring either stronger persistence or observed performance degradation for retraining.

7.2. Predictive Adaptation Is Not Policy Optimization

The higher normalized cost of DARC-OR is not hidden as an inconvenient secondary metric. It is a central result. The experiment demonstrates that a detector can improve AP, recall, and Brier score without improving an exposure-weighted decision objective. This measured trade-off reinforces the useful architectural separation in XOR-EW between ranking/calibration evidence and policy optimization [4]. In practice, the institution should validate the economic coefficients and investigation capacity independently, then replay alternative policies on frozen score histories before changing the model. DARC-OR should be viewed as a drift-control layer feeding such a policy process, not as an end-to-end claim of minimum financial loss.

7.3. Recurrence Reveals the Missing Memory Layer

The recurrence result is equally important for innovation. Short-memory retraining is an efficient response to abrupt novelty, but it can destroy a representation that later becomes useful again. DARC-OR performs well in B/C precisely because it forgets A aggressively; that same design causes the large AP deficit when A returns. This failure suggests a next-generation architecture with a versioned concept library. Before fitting a new model, the controller could compare the current score/feature signature against stored regimes and activate a validated historical champion when similarity and safety criteria are satisfied. Such a memory layer should be tested with contamination controls so that an attacker cannot deliberately induce unsafe model reuse.

7.4. Deployment Blueprint

A controlled deployment should begin in shadow mode.

The institution would log raw score distributions, delayed outcomes, JSD, balanced-loss EWMA, action recommendations, selected memory windows, capacity usage, and realized outcomes while leaving the incumbent decision policy unchanged. Approval gates would be separate for calibration updates, model retraining, and alert-policy changes. A recalibration may be authorized under a lighter change process than a new model, whereas an adaptive retrain should trigger validation of performance by business segment, stability, operational latency, and reason-code continuity. The current paper does not evaluate explanations; for deployment, DARC-OR should be combined with an institution-specific explanation and audit layer rather than treating drift statistics as explanations of individual events.

8. Threats To Validity

First, the benchmark is synthetic. That choice provides exact drift timing and avoids inventing labels, but it cannot reproduce the full dependence structure, control environment, loss taxonomy, adversarial behavior, or reporting delays of a financial institution. Transaction fraud is only a narrow operational-loss proxy. The study therefore makes no claim of production readiness or universal superiority.

Second, the five seeds are independent executions of one generator family, not five institutions. Their sample size is sufficient to show repeatability of the controlled mechanisms but too small for strong inferential statistics. Third, detector thresholds (JSD floor 0.004, robust multiplier 5, EWMA factor 0.35, loss ratio 1.35, persistence counts, and cooldown) are engineering choices fixed before the final comparisons; they are not universally optimal. Fourth, the event-window delay metric can associate an action with a known change even when the controller state was influenced by prior deterioration, as seen at day 70.

Fifth, the normalized exposure cost is a dimensionless scenario. It is useful for controlled comparison but is not a bank loss estimate. Different coefficients or review capacities can reorder methods. Sixth, DARC-OR lacks a recurring-concept memory mechanism and therefore exhibits severe forgetting when A returns. Finally, fairness, protected-attribute behavior, investigator workload time, cybersecurity resilience, adversarial manipulation, explanation stability, and real-world delayed-label distributions are outside the present experiment and require separate validation before operational use.

9. Conclusion

DARC-OR introduces a dual-signal drift-response layer for capacity-constrained operational-risk early warning. Immediate JSD monitoring supplies label-free evidence of score-distribution change; delayed balanced-loss EWMA supplies performance evidence; and a tiered controller chooses recalibration or recency-weighted 14/30-day retraining. On five executed controlled streams totaling 720,000 transactions, DARC-OR achieved the highest mean AP (0.1719), improved recall relative to delayed-loss retraining, and reduced Brier score relative to a static model. These gains were strongest after novel B and C regimes. At the same time, DARC-OR did not minimize the exposure scenario cost and performed substantially worse when the original A regime recurred. The combined result is more useful than a single headline metric: concept-drift adaptation can improve warning quality, but action economics and

concept memory remain independent governance problems. Future work should evaluate a validated concept-memory library, realistic stochastic label delays, institution-calibrated costs, and multi-month operational-loss data under shadow deployment.

References

- [1] Zhang, H. (2024). Graph-aware multi-task learning for early prediction of timing violations and routing congestion in digital IC physical design. *World Journal of Information Technology*, 2(1), 13–20.
- [2] Wang, J., Tan, Y., Jiang, B., Wu, B., & Liu, W. (2025). Dynamic marketing uplift modeling: A symmetry-preserving framework integrating causal forests with deep reinforcement learning for personalized intervention strategies. *Symmetry*, 17(4), 610.
- [3] Zhao, X., Liu, J., Wang, Y., & Wang, J. (2026). CryptoMamba-SSM: Linear complexity state space models for cryptocurrency volatility prediction. *IEEE Open Journal of the Computer Society*, 7, 226–243.
- [4] Wang, J., Liu, J., Zheng, W., & Ge, Y. (2025). Temporal heterogeneous graph contrastive learning for fraud detection in credit card transactions. *IEEE Access*.
- [5] Sun, T., Yang, J., Li, J., Chen, J., Liu, M., Fan, L., & Wang, X. (2024). Enhancing auto insurance risk evaluation with transformer and SHAP. *IEEE Access*, 12, 116546–116557.
- [6] Zhang, H. (2025). Reinforcement learning approaches for layout optimization in electronic design automation with electromagnetic compatibility constraints. *Frontiers in Robotics and Automation*, 2(2), 77–93.
- [7] Han, X., Yang, Y., Chen, J., Wang, M., & Zhou, M. (2025). Symmetry-aware credit risk modeling: A deep learning framework exploiting financial data balance and invariance. *Symmetry*, 17(3), 341.
- [8] Chen, J., Cui, Y., Zhang, X., Yang, J., & Zhou, M. (2024). Temporal convolutional network for carbon tax projection: A data-driven approach. *Applied Sciences*, 14(20), 9213.
- [9] Sun, T., Wang, M., & Chen, J. (2025). Leveraging machine learning for tax fraud detection and risk scoring in corporate filings. *Asian Business Research Journal*, 10(11), 1–13.
- [10] Chen, J., Wang, M., & Sun, T. (2025). Intelligent tax systems and the role of natural language processing in regulatory interpretation. *American Journal of Machine Learning*, 6(4), 74–94.
- [11] Chen, Z., Liu, J., & Chen, J. (2025). Machine learning methods for financial forecasting in enterprise planning: Transitioning from rule-based models to predictive analytics. *Frontiers in Artificial Intelligence Research*, 2(3), 541–564.
- [12] Chen, J., Liu, J., Liang, Y., & Zhou, M. (2026). KE-MLLM: A knowledge-enhanced multi-sensor learning framework for explainable fake review detection. *Applied Sciences*, 16(6), 2909.
- [13] Yue, X., Yang, J., & Liu, W. (2026). FraudDebate-Agent: A multi-agent LLM framework with an evidence-based debate mechanism for financial statement fraud detection. *Mathematics*, 14(15), 2695.
- [14] Zou, J., Qin, Y., & Rhee, M. (2024). Beyond model scale: A task-aware reliability benchmark for large language models in clinical decision support. *Innovation and Technology Studies*, 1(1), 39–47.
- [15] Yue, X., & Zhu, R. (2024). An explainable hybrid machine learning framework for cash-flow-conditioned multi-horizon liquidity risk early warning: An SME-oriented benchmark study. *World Journal of Management Science*, 2(1), 63–72.

- [16] Zhang, C., & Wang, Y. (2024). A data-driven multi-objective production scheduling framework for battery pack manufacturing under material arrival uncertainty. *AI and Data Science Journal*, 1(1), 86–93.
- [17] Zi, B. (2024). Cloud-native distributed systems for real-time payment intelligence. *AI and Data Science Journal*, 1(1), 51–56.
- [18] Zhang, S., & Qiu, L. (2023). An interpretable, uncertainty-aware framework for bridge deterioration prediction and risk-based maintenance prioritization. *Academic Journal of Architecture and Civil Engineering*, 1(2), 21–28.
- [19] Jin, J., Xing, S., Ji, E., & Liu, W. (2025). Xgate: Explainable reinforcement learning for transparent and trustworthy API traffic management in IoT sensor networks. *Sensors*, 25(7), 2183.
- [20] Xing, S., Wang, Y., & Liu, W. (2025). Multi-dimensional anomaly detection and fault localization in microservice architectures: A dual-channel deep learning approach with causal inference for intelligent sensing. *Sensors*, 25(11), 3396.
- [21] Xing, S., & Wang, Y. (2025). Proactive data placement in heterogeneous storage systems via predictive multi-objective reinforcement learning. *IEEE Access*, 13, 117986–117998.
- [22] Ji, E., Wang, Y., Xing, S., & Jin, J. (2025). Hierarchical reinforcement learning for energy-efficient API traffic optimization in large-scale advertising systems. *IEEE Access*.
- [23] Xing, S., & Wang, Y. (2025). Cross-modal attention networks for multi-modal anomaly detection in system software. *IEEE Open Journal of the Computer Society*.
- [24] Liu, Y., Ren, S., Wang, X., & Zhou, M. (2024). Temporal logical attention network for log-based anomaly detection in distributed systems. *Sensors*, 24(24), 7949.
- [25] Ren, S., Jin, J., Niu, G., & Liu, Y. (2025). ARCS: Adaptive reinforcement learning framework for automated cybersecurity incident response strategy optimization. *Applied Sciences*, 15(2), 951.
- [26] Li, P., Ren, S., Zhang, Q., Wang, X., & Liu, Y. (2024). Think4SCND: Reinforcement learning with thinking model for dynamic supply chain network design. *IEEE Access*, 12, 195974–195985.
- [27] Hu, X., Zhao, X., Wang, J., & Yang, Y. (2025). Information-theoretic multi-scale geometric pre-training for enhanced molecular property prediction. *PLOS One*, 20(10), e0332640, https://doi.org/10.1371/journal.pone.0332640[https://doi.org/10.1371/journal.pone.0332640].
- [28] Wang, M., Zhang, X., Yang, Y., & Wang, J. (2025). Explainable machine learning in risk management: Balancing accuracy and interpretability. *Journal of Financial Risk Management*, 14(3), 185–198.
- [29] Yang, J., Li, P., Cui, Y., Han, X., & Zhou, M. (2025). Multi-sensor temporal fusion transformer for stock performance prediction: An adaptive Sharpe ratio approach. *Sensors*, 25(3), 976.
- [30] Zhao, X., Sun, T., Ren, S., Yang, J., & Liu, Y. (2025). RAG-based AI agents for enterprise software development: Implementation patterns and production deployment. *Frontiers in Artificial Intelligence Research*, 2(3), 501–520.
- [31] Guo, Z., Chen, T., & Shang, W. (2022). Class-conditional intermediate-domain adversarial adaptation for cross-device acoustic scene classification. *Journal of Computer Science and Electrical Engineering*, 4(1), 26–35.
- [32] Wang, Y., Bian, M., & Lin, Y. (2022). An explainable temporal graph neural network for disruption prediction and risk propagation in multi-tier supply chains. *Journal of Computer Science and Electrical Engineering*, 4(1), 17–25.
- [33] Zhao, W., & Wang, B. (2022). CAPA: A prediction-driven autoscaling framework for SLO-aware cloud-native machine learning inference. *Journal of Computer Science and Electrical Engineering*, 4(1), 8–16.
- [34] Jiao, Y., & Liang, Y. (2024). An integrated process-mining and explainable machine-learning framework for bottleneck and internal-control risk detection in the multi-entity financial close. *Social Science and Management*, 1(3), 96–107.