

Lightweight Anomaly Detection over Digital-Twin State Streams

F. Farah and G. Gupta

2023 (synthetic sample)

Synthetic sample paper. Written solely as a sample source for the chi-tragupta documentation’s worked examples; the methods are textbook and the results are invented for illustration. Public domain (CC0).

Motivation

A digital twin’s state stream is a natural place to detect anomalies: it is already cleaned, aligned and semantically labelled, which is most of the work in industrial anomaly detection. The question this paper examines is how little model is enough once the twin has done that work.

Three detectors

We compare three detectors of increasing cost on the same synthetic state streams. The first is a **rolling z-score** per signal: a value more than four standard deviations from its trailing one-hour mean is flagged. The second is an **isolation forest** over ten-signal windows, retrained nightly. The third is a small **autoencoder** whose reconstruction error is thresholded at the ninety-ninth percentile of a clean week.

Results

On streams with injected point faults, all three detectors exceeded 0.95 recall; the z-score’s precision (0.71) trailed the forest (0.88) and the autoencoder (0.90) because it fired on legitimate setpoint changes it had no context to recognise. On injected drift faults the ordering reversed at the tail: the autoencoder, retrained on a window that had begun to drift, absorbed the fault into its notion of normal and missed the slowest ramp entirely, while the humble z-score – anchored to a longer baseline – flagged it after nine hours.

The cost gap is larger than the accuracy gap. The z-score runs in constant memory per signal; the forest’s nightly retrain took eleven minutes on our synthetic plant;

the autoencoder required a GPU hour per week and a human decision about when a retraining baseline is clean.

Recommendation

Layer, do not choose. Run the z-score everywhere as the incorruptible baseline nobody retrains, and add a learned detector only on the signals whose false-positive cost justifies its upkeep. The layered configuration matched the autoencoder's precision on point faults while retaining the z-score's immunity to baseline corruption on drift – the failure mode that matters most, precisely because it is the one a learned model conceals.