

A decision engine, not a scorecard. The difference, stated plainly.

Written for buyers who have seen a lot of dashboards and want to know what is actually different here.

Start with the concession: the scoring layer is a weighted scorecard. Signals in, weights applied, a composite score out.

That part is not novel and we will not pretend otherwise. What makes SignalOS an engine is what is built around the score – and those three things are where the value sits.

Where a scorecard stops and an engine begins

A SCORECARD	SIGNALOS
Reports that a channel scored 0.6. A person reads it and decides what that means.	Produces an instruction. “Accelerate, +5% to +10%, capped at 15% week-on-week, because demand and conversion are strengthening on high-reliability data.” Remove the human and a scorecard yields nothing; this yields an executable decision.
Has no veto and no null state. Whatever the inputs, it produces a number.	Can override its own score, and can refuse to run. A margin failure signed by Finance forces a cut at any score. A missing sign-off blocks the recommendation entirely rather than defaulting past it.
Never finds out whether it was right.	Closes the loop. Decisions are tested against outcomes each quarter, and that evidence drives recalibration of the engine itself.

How the signals are weighted

Eight signals are graded each week, per channel, as Up, Neutral or Down – Demand, Cost, Conversion, Margin, Competition, Retention, Event and Creative. Four steps turn that into a decision:

01	Direction is applied. Raw movement is converted into commercial favourability. Rising costs and intensifying competition count <i>against</i> a channel, not for it – an inversion that is easy to get wrong and expensive when you do.
02	Reliability is weighted. Each signal carries a confidence grade set by the quality of its source that week. High-confidence data counts double; weak data counts half; missing data is excluded from the calculation entirely and flagged on the recommendation. Nothing is silently assumed.
03	The score is normalised. The weighted total is divided by the total weight applied, producing a bounded score. This matters more than it sounds: it means the decision thresholds behave identically in a week when your data is strong and a week when half your feeds are degraded. Poor data lowers confidence, it does not distort the verdict.
04	Gates are applied last. The margin gate, the demand and conversion condition, and the week-on-week movement cap are evaluated after the score, and can override it.

How channel and market differences are handled

Per channel, weekly	Every signal and every reliability grade is set at channel level. Affiliates, PPC, programmatic and paid social are each scored on their own evidence, and each receives its own decision and its own rationale.
Per market	Commercial intent is set per market and owned by you: allowable CPA, payback window, margin gate definition, risk mode, creation and harvest split, and retention targets. Two markets running the same signals can legitimately reach different decisions because their intent differs.
Per implementation	Signal derivation rules are calibrated to your data during the diagnostic and shadow period – what constitutes a meaningful cost movement in a mature market is not what it means in a launch market.
Per calendar	Event tiering reflects the local sporting and promotional calendar, easing the acceleration threshold when a genuine demand event is in the window.
On the roadmap	Threshold profiles set per channel as well as per market. Today thresholds are calibrated at market level and applied consistently across channels within it; per-channel profiles are the next planned extension of the engine.

How recommendations evolve based on results

This is the question that separates a governed engine from a static model, and the answer has two halves – one principled, one mechanical.

- 01 **The principle: no silent self-tuning**

The engine does not quietly adjust its own weights between cycles. It could, technically. It does not, deliberately – because a model that rewrites its own rules cannot be replayed, and a decision that cannot be replayed cannot be audited. Determinism is the property the whole framework exists to protect.
- 02 **The mechanism: a governed quarterly cycle**

Recalibration happens at a defined point, on evidence, as a documented act with a named owner and a signed change record showing what moved, why, and what it would have changed last quarter.

The quarterly loop

1 • RECORD Every recommendation, override and executed change is logged as it happens, with its rationale and approver.	2 • VERIFY Decisions are tested against what actually happened: acquisition cost against allowable, payback against window, margin against expectation.	3 • ATTRIBUTE Accuracy is assessed signal by signal. When Demand read Up, did those channels deliver? Signals that mislead lose weight; signals that prove reliable gain it.	4 • RECALIBRATE Thresholds and derivation rules are adjusted on that evidence, signed off, and issued as a dated calibration record.
--	--	---	---

What it costs you to run

A fair question, and one we have historically undersold. SignalOS is an operating discipline with clearly separated roles, and the decision-maker's share of it is small.

ROLE	WHAT THEY DO EACH WEEK	TIME
Budget owner (often the CMO)	Reviews the week's recommendations, approves, adjusts within range, or rejects with a recorded reason	15-30 min
Finance signatory	Sets the margin gate to pass or fail per product-market	~5 min
Engine operator	Enters signals and reliability grades, records decisions, confirms changes landed	~45 min
Everyone else	Nothing	–

The decision-maker never touches the model. They see a one-page readout and make a call on it. And where operators automate signal collection – which the published integration specification supports – the operator's entry work reduces to exception handling.

What is deliberately not in this document

The specific thresholds, the signal derivation rules that convert your platform data into Up, Neutral or Down, and the weighting profile applied to your markets. Those are calibrated to each client during implementation and constitute the intellectual property of the framework. Everything above describes the mechanism; the calibration is what we are engaged to get right, and it is documented in your own calibration record from day one.

THE SHORT VERSION

The score is a scorecard. The engine is the gates that can override it, the bounded and executable output it produces, and the quarterly loop that tests it against reality and recalibrates on the evidence. Take away any one of those three and what remains is a very good dashboard.