

The detection evidence record

A P–F interval is a claim about one failure mode, one functional boundary, one measurement method, one operating context and one decision process. Six fields hold the claim; four questions test it. A number carried in from another site fills none of them.

RECORD Six fields, filled before an interval is discussed

- 01 Function and failure mode.** Required function, performance boundary, failure mechanism, consequence context. Written per mode, so an asset-level interval cannot hide unlike modes.
- 02 Measurement contract.** Technique, sensor or sample location, operating state, acquisition settings, data-quality checks — what “detectable” means in practice here.
- 03 Detection criterion.** Baseline, criterion, persistence or confirmation rule, and whether an analyst or a system decides.
- 04 Warning evidence.** Relevant field or test observations, population and conditions, spread, censoring, missing data, and the limits on transfer.
- 05 Response path.** Alert review, diagnosis, decision authority, parts, access, isolation, work window, execution — how much of the warning is consumed after detection.
- 06 Governance.** Consequence class, applicable law and regulation, OEM or controlled-procedure requirements, owner, approvers, review triggers.

ASK Four questions kept separate, so missing information stays visible

- **What warning evidence exists?**
Population and conditions, exposure basis, observed spread, censored or missing observations, transfer limits, and the reasons the evidence may no longer apply. A mean, one anecdote or a borrowed asset-class number is not a defensible lower bound.
- **When can a valid measurement actually occur?**
Operating-state opportunity, access, missed or invalid readings, analysis, confirmation, escalation — and failure of the monitoring system itself.
- **What does the response path demand?**
Diagnosis, decision authority, parts, access, permits and isolation, work-window alignment, repair and return-to-service controls, reviewed under credible adverse conditions rather than the best case.
- **What uncertainty and consequence controls apply?**
Mandated and OEM tasks, protective functions, safe-state and shutdown rules, independent safeguards, change control, decision authority, and the site-approved basis for any added conservatism.

THE BOUNDARY THIS CARD CARRIES

If the warning evidence, measurement performance, response path or approval basis is insufficient, the result is **NOT DEMONSTRATED** — and a plausible-looking timeline is not permission to change an interval. This card sets no interval, threshold or fraction, and authorises no lubrication, inspection, proof-test or protective-function change.

Ordering the queue

Two questions per machine — how much does it hurt when this fails, and how often does that happen — answered against written scales and applied to every asset with the same yardstick. What comes out is an ordered queue, not a risk value.

RECORD What every scored row has to carry

- 01 The slot, not the machine.** Criticality belongs to the functional location. The duty position is scored; the identical rebuilt spare in the workshop is not — yet.
- 02 Consequence, criterion by criterion.** Safety and environment, production, quality, cost of the failure event, spares and repair response — each scored against the worst *credible* event, with the reason written in the cell.
- 03 Likelihood, with its source.** The events-per-year anchor, plus where the score came from: work orders, class history, or structured judgement. Flag the judgement so data can replace it later.
- 04 The combination rule.** Which criteria, what scale, how aggregated, how weighted, and where the bands cut — written on the front of the workshop pack before scoring starts.
- 05 Overrides.** Statutory inspection and test, protective and safety-instrumented functions, environmental permit conditions, insurer and OEM conditions, product-quality, regulatory and contractual obligations, flagged on the register entry.
- 06 Date, room and assumptions.** Who scored it, when, and which assumptions — duty, redundancy, consequence, repair route — the score depends on, because it decays as those change.

ASK Before the room signs the list

- **Was the rule fixed before anyone could see who wins?**
Each aggregation and weighting choice produces a different order, and none is more correct in the abstract. What makes the ranking defensible is that nobody adjusted the rule once the result was visible.
- **Who was in the room?**
Maintenance alone overweights repair cost; operations alone overweights last month's pain. The score needs maintenance, operations, a process engineer, and someone who can speak for safety and cost.
- **Worst credible, or worst imaginable?**
Meteor-strike scoring inflates everything to the top of the scale and destroys the ranking's resolution. Say which event you scored against.
- **Has the standby been demonstrated, or is it a rumour?**
A duty/standby pair with a tested changeover earns a low production consequence. An untested standby scores as if absent.
- **Which scores rest on evidence nobody can produce?**
Where the basis is missing or disputed, record **NOT DEMONSTRATED** and keep the entry open rather than settling on a confident number.

THE BOUNDARY THIS CARD CARRIES

These are ordinal scores. An asset ahead of another sits earlier in the queue; it is not “twice as critical”, and the product is not a risk figure in any measurable currency. No band shortens, defers or deletes a statutory, protective, environmental, quality or contractual obligation, and this card proposes no scale, weighting or band cut of its own.

Four words, one sentence

“It broke” aims no sensor. Written as mode, mechanism, cause and effect, the same event tells you where to look, what might be observable in the data, what to fix so it stops recurring, and why anyone should pay for the watching.

RECORD One mode, one row — four words and the two things that are not modes

- 01 Functional failure.** What the operating context requires, and the boundary at which the asset no longer meets it. The equipment may still be running.
- 02 Failure mode.** The way that loss of function occurs, named at maintainable-item level. The asset’s name is not a mode.
- 03 Failure mechanism.** The physical or chemical process doing the damage — fatigue, abrasive wear, corrosion, erosion, fouling, insulation breakdown, fretting. Named, not implied.
- 04 Failure cause.** Why the mechanism started or accelerated. Kept separate from the mechanism: causes are what root-cause work removes, mechanisms are what sensors watch while the causes persist.
- 05 Effect, at three levels.** Local, next level, end effect — with severity judged at the end effect, in this plant, this duty and this redundancy state.
- 06 Symptom — not a mode.** “Runs hot”, “trips on overload”, “noisy at the drive end”: useful as evidence pointers, useless as rows, because several modes produce each of them.
- 07 Detectable condition — narrower still.** Not “vibration”, but a defined change in a named component of the signal, measured at a stated location and load, against a comparator taken in the same state.

ASK Run each row past these before it enters the worksheet

- **Read right to left, does the row make a sentence?**
Cause, then mechanism, then mode, then effect, as one readable clause. If it does not read as a sentence, a word is missing or two are blurred.
- **Does the row have a mechanism?**
A row without one cannot be monitored, only mourned — the mechanism is the physics a measurement would have to go looking for.
- **Are two modes merged because they share a component?**
“Bearing seizes” and “bearing runs hot from misalignment” are different modes with different evidence. Merging them blinds the measurement choice.
- **Is severity scored at the end effect?**
The identical component failure in a standby with a demonstrated changeover has a trivial end effect. Modes travel between plants; effects do not.

THE BOUNDARY THIS CARD CARRIES

Read left to right, the row is a monitoring *question* — not yet a monitoring specification. Whether anything is detectable early enough to act on is a separate question, answered per mode and per context by evidence, never by the grammar. Where that evidence is absent the row stands at **NOT DEMONSTRATED**.

The measurement-evidence record

A sensor can work exactly as designed and still support no defensible maintenance decision. Nine fields decide whether a condition signal is evidence about a named failure mechanism or a datasheet feature that happens to be trending.

RECORD Nine fields, written before anyone quotes a price

- 01 Required function and failure mode.** Function, functional boundary, mode, mechanism, operating context, consequence context — so an asset name cannot hide unlike mechanisms.
- 02 Candidate signal.** The physical quantity or observation, and the reason it may change as the mechanism develops.
- 03 Measurement contract.** Technique, location, direction or orientation, operating state, acquisition method, settings, and the conditions that make an observation valid.
- 04 Data-quality controls.** Calibration or verification status, mounting or sampling integrity, range, resolution, clipping or contamination checks, timestamps, context completeness.
- 05 Detection and confirmation rule.** Baseline or comparator, criterion, persistence requirement, confirmation method, and the competent reviewer — so one number or image cannot become a diagnosis.
- 06 Warning evidence.** Relevant observations, exposure basis, variability, misses, invalid observations, and the limits on transfer.
- 07 Failure and uncertainty modes.** Confounders, blind spots, shared dependencies, monitoring-system failure, state-dependent misses, alternative explanations.
- 08 Response path.** Review, diagnosis, decision authority, parts, access, isolation, work window, execution, return-to-service controls.
- 09 Governance and status.** Applicable law, OEM requirements, controlled procedures, protective functions, decision owner, approvers, review triggers, evidence status.

ASK The record produces one of three states — and one trap

- **SUPPORTED FOR THIS CONTEXT**
Relevant evidence exists within the stated conditions and transfer limits. An evidence classification, not permission to alter an operating or maintenance requirement.
- **CONTROLLED TRIAL**
Plausible, with local measurement performance, warning evidence, confirmation or response performance still to be established under approved change control.
- **NOT DEMONSTRATED**
The physical link, measurement contract, data quality, warning evidence, confirmation, response path or approval basis is incomplete. A legitimate result, not a failure of nerve.
- **Has it been quiet, or has it been blind?**
Silence does not promote a signal. A quiet trend can equally mean an insensitive measurement, an invalid operating state, a missed observation, or a channel that died last Tuesday.

THE BOUNDARY THIS CARD CARRIES

Technical suitability and economics are separate judgements, and neither validates the other. This card selects no monitoring technology, establishes no diagnosis, sets no alarm or protection value, prescribes no inspection or route interval, and authorises no installation, testing, maintenance or continued operation.

Seven fields, five rungs

The consumer is a planner or technician who was not watching the screen and has eleven other things due. An alert is a complete utterance or it is noise with a timestamp — and the finding behind it stands on exactly one rung.

RECORD Seven fields, on one screen

- 01 Identity.** Asset and maintainable item, by hierarchy address.
- 02 Finding.** Which feature moved, against which basis, in which operating state.
- 03 Meaning.** The suspected failure mode, named from the library. A candidate that cannot fill this field is a threshold event for an analyst's queue, not a dispatch.
- 04 Severity.** The triage band, from consequence and urgency.
- 05 Confidence.** How sure, in calibrated words.
- 06 Time.** The window and act-by date, however coarse.
- 07 Action and owner.** The recommended next step and the named role it routes to. A message with no owner is a dashboard nobody owns, refused at birth.

ASK Which rung does this finding stand on?

- 01 Anomaly**
A feature moved against its declared basis, in a valid operating state. Supports a look: triage, a tightened watch, a request for a second measurement.
- 02 Suspected failure mode**
The evidence fits a named mode from the library, rivals still open. Supports an investigation with a stated question and an evidence list.
- 03 Confirmed defect**
Corroborated by an independent witness or direct inspection, with the expected as-found written down. Supports a work recommendation.
- 04 Work recommendation**
A proposed scope, resources and act-by date drawn from the site's own action library. Supports planning, scheduling and procurement review.
- 05 Authorised operating decision**
Run, restrict, stop, restart, return to service. Not on this card, and not in the monitoring system.

THE BOUNDARY THIS CARD CARRIES

A monitoring system can climb to the fourth rung on its own evidence. The fifth is not a stronger alert — it is a different kind of decision, made by the people who hold operating, process-safety and integrity authority and answer for the consequence. Every possible action leaves this card as an escalation question with its evidence requirement attached, never as an instruction to run, restrict or stop.

The registration sheet

A verdict negotiated after the evidence arrives is not a verdict but a settlement. Fill this page in before day zero, file it where it cannot be edited without a trace, and let the review read the page instead of the room.

RECORD Six lines, filed before the first phase begins

- 01 Primary outcome.** What would count as success, stated in the window's own terms, priced from your own per-event table and agreed with finance — the strongest claim a window of that length can carry.
- 02 Declared undecidable, in advance.** The questions the window is too short to settle, registered as **NOT DEMONSTRATED** whatever happens, and carried to the longer horizon agreed with finance.
- 03 Stop line.** The avoided share below which the case fails *at your horizon*: your committed spend divided by your horizon multiplied by your own annual baseline. Computed from your own numbers, never borrowed.
- 04 Secondary outcomes.** An alert-precision floor and a disposition-latency contract, both set from your own alert volumes and staffed hours; and no finding living outside the system of record.
- 05 Explicitly not outcomes.** Team enthusiasm, demo quality, dashboard aesthetics, and any measure invented after day zero.
- 06 The date and the jury.** The registered end date, the review tier it is heard at, and who is in the room — sponsor and finance included.

ASK Four rules that keep the goalposts bolted down

- **Is the analysis registered, not just the measure?**
Write down how each number will be computed — the counterfactual per event, priced from your own table — or the computation quietly becomes the goalpost.
- **Are the readings pre-committed?**
SUPPORTED FOR THIS CONTEXT for the bounded claim the evidence actually carries, **NOT DEMONSTRATED** where it does not, plus stop, or restructure once against one named defect with one further window.
- **Does the page say what the window cannot settle?**
A question registered as undecidable in advance cannot be quietly answered in month three.
- **Did the length come from the failure or the diary?**
The window is set by how often the failure you are trying to move actually occurs, and by the cadence at which evidence about it arrives. Then the clock starts at first data, not at kickoff, and nothing new enters the proving phase.

THE BOUNDARY THIS CARD CARRIES

NOT DEMONSTRATED is a verdict in its own right, not a failure of nerve, and a programme that says it out loud at the first review is the one whose second review gets believed. None of the verdicts approves a control, architecture, supplier, maintenance action or operating decision, and this card supplies no window length, acceptance value or stop threshold of its own.