

TRAINING MODULE

Engineering Review Board (ERB)

Name the real constraint, or admit there isn't one

A working module for engineers with 1–2 years on the floor. By the end you will be able to build a station fingerprint from real data, and run a five-round Socratic debate that doesn't let a habitual step hide behind “that's just how it's done.”

Grounded in the Socratic method and TWI Job Methods' “question every detail” discipline, formalized into a five-round structure

What You Will Be Able to Do After This Module

- 1 Explain how ERB differs from Judgement Review — evidence-triggered debate vs. a single captured idea.
- 2 Build a station fingerprint from element text, the same way the keyword lexicons do.
- 3 Explain the five-round escalation structure, and what each round is actually forcing the engineer to do.
- 4 Run a real Round 1–5 debate on a genuine upstream-defect pattern, by hand.
- 5 Explain exactly what the customer's own IE can and cannot edit in ERB_LIBRARY.xlsx — and why that ceiling exists.

A Captured Idea Is Not the Same as a Forced Debate

You've already met Engineering Judgement Review — a single floor idea, honestly logged. ERB is a different tool, for a different trigger:

JUDGEMENT REVIEW

An engineer notices something and logs it. One idea, one honest review, one decision.

ENGINEERING REVIEW BOARD

A measurable FINGERPRINT in the data triggers the debate automatically — nobody has to notice anything. Then five rounds refuse to let a habitual answer stand unchallenged.

Judgement Review captures what someone already saw. ERB goes looking for patterns nobody necessarily noticed — and then won't let the answer be “that's just how it's done.”

Debates Are Triggered by Evidence, Not Opinion

Every station gets a FINGERPRINT — 18+ measurable fields, computed from real element text and timing. A debate only opens when the fingerprint actually matches a trigger condition:

has_wait	has_setup	has_transport	has_inventory
has_measure	has_rotate	has_machine_trigger	has_decision
has_upstream_trigger	micro_ct_flag	repeat_count	+ 8 more

A debate never opens because a station “feels wasteful.” It opens because the fingerprint matches a real, stated trigger condition in ERB_LIBRARY.xlsx — evidence-bound from Round 1.

WHERE THIS METHOD COMES FROM

You Already Met This Discipline — Now It Escalates

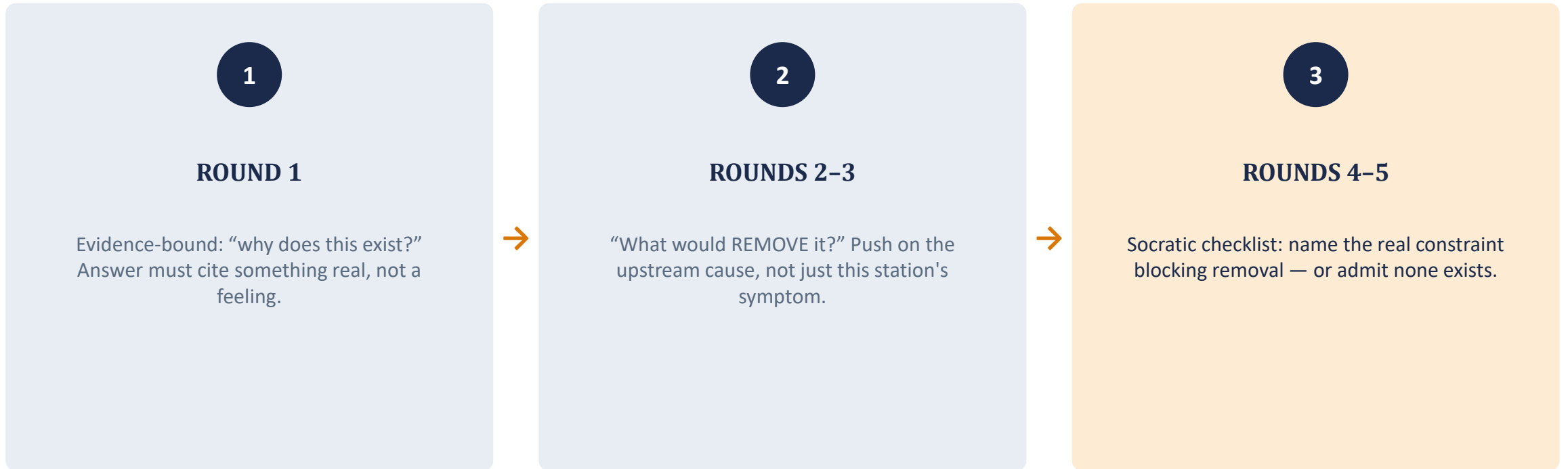
TWI Job Methods asks “why is it necessary” once. The Socratic method's real power is refusing to accept the FIRST answer as final — asking a follow-up that tests whether the first answer was a reason or just a habit. ERB formalizes that refusal into five rounds: Round 1 demands evidence for existence, Rounds 2–3 push on what removal would take, and Rounds 4–5 force a checklist that ends in exactly one of two places — a named, real constraint, or an admission that none exists.

“We've always done it this way” is not an answer this structure lets survive.

SOCRATIC METHOD

+ TWI Job Methods
“Question Every Detail”

Five Rounds, Each With a Different Job



Every triggered activity walks this same arc — the rest of this module runs one all the way through.

Three Steps, Every Time

1

COMPUTE FINGERPRINT

Keyword lexicons scan every element's text and timing — the measurable facts, nothing inferred.

2

MATCH AGAINST ERB_LIBRARY.xlsx

Check which trigger conditions the fingerprint actually satisfies.

3

RENDER THE MATCHING ROUNDS

Only the rounds whose trigger fired get shown — never a generic debate.

1

STEP 1

Build the Fingerprint

Every keyword lexicon runs against every element's real text — nothing about the fingerprint is guessed.

MICRO_CT_SECONDS = 2.0s and REPEAT_MIN_COUNT = 3 are the two thresholds that turn raw text into a flag.

WORKED ELEMENT

“Operator waits for machine cycle, then deflashes part edge with a hand file.” — Observed CT: 1.5s

Field	Match	Value
has_wait	“wait”	TRUE
has_upstream_trigger	“deflash”	TRUE
micro_ct_flag	1.5s < 2.0s threshold	TRUE
repeat_count	same pattern at 4 stations	4 (≥ 3)

Four fingerprint fields fired TRUE on one 1.5-second element — that combination is exactly what a trigger condition looks for.

2

STEP 2

Match Against the Library

The customer's own IE can edit this exact condition next year — wording, thresholds, even add Round 6 — with zero code changes.

What they CANNOT do: invent a brand-new fingerprint field the Python side never computed.

THE TRIGGER CONDITION, FROM ERB_LIBRARY.xlsx

**IF `has_upstream_trigger == TRUE` AND `repeat_count >= 3`
THEN open Board: “Upstream Root-Cause Debate”**

CHECKING OUR FINGERPRINT AGAINST IT

`has_upstream_trigger == TRUE` — matches (“deflash”)

`repeat_count (4) >= 3` — matches

BOTH conditions true → Board OPENS: “Upstream Root-Cause Debate”

This same match happened independently at all 4 stations sharing this pattern — the `repeat_count` itself came from noticing that.

“Upstream Root-Cause Debate” — Worked Full

R1

This deflash step appears at 4 stations, combined ~6s/cycle across the line. What upstream process actually creates the flash being removed here?

R2–3

If the molding parameter causing this flash were corrected at the source, could all 4 downstream deflash steps be eliminated? What would that correction cost, vs. the combined removal cost today?

R4–5

Is there a genuine technical reason flash cannot be prevented at the mold — material, tooling wear, cycle-time trade-off? Name it specifically, or confirm no such constraint exists.

The engineer must answer R4–5 with a **NAMED** constraint or an admission — “we never checked” is itself a valid, honest Round 5 answer that opens the real next step.

What the Customer's Own IE Can — and Can't — Edit

There will be no ongoing vendor relationship after delivery. Both of these must be true at handover:

CAN EDIT, NO CODE NEEDED

Question wording, trigger thresholds (e.g. `repeat_count ≥ 3` → `≥ 4`), even add a brand-new Board or Round 6 — all live in ERB_LIBRARY.xlsx.

CANNOT, WITHOUT A CODE CHANGE

Invent a brand-new fingerprint field — e.g. “walking distance in metres” — if nothing in the source data measures that yet. The fingerprint is built wide (18+ fields) to keep this ceiling as high as possible.

This is a stated, honest limitation — not hidden. Knowing the ceiling exists is exactly what lets you set the customer's expectations correctly at handover.

Five Mistakes That Give Junior IEs Away



Accepting “that's just how it's done” as a valid Round 4–5 answer — it isn't a named constraint, and the round isn't finished.



Opening a debate on a station that doesn't actually match a trigger condition — evidence-bound means the fingerprint has to fire, not a hunch.



Trying to add a metric the fingerprint never computes by editing ERB_LIBRARY.xlsx alone — that needs a real code change, not a spreadsheet edit.



Treating Round 1's answer as the final word — the whole point of Rounds 2–5 is to keep pushing past the first, comfortable answer.



Losing the accumulated debate history between sessions by skipping the Export/Load Answers step — the discussion should build, not restart.

What generate_erb_v9.py Does For You

A static, fully offline HTML debate — no server, no silent data capture, and a real memory across sessions.

1

Computes every station's fingerprint

All 18+ fields, from real element text and timing — the Slide 8 logic, run automatically across the whole line.

2

Matches against ERB_LIBRARY.xlsx

Checks every trigger condition and opens only the Boards that actually fire — never a generic, one-size debate.

3

Renders the matching rounds, in order

Evidence-bound Round 1 through the Socratic Round 4–5 checklist — exactly the structure from Slide 10.

4

Export / Load Answers

A single JSON file the browser downloads and re-opens — the debate can be revisited and accumulated across multiple review sessions, with nothing saved silently anywhere else.

KEY TAKEAWAYS

Recap — Say This Back In Your Own Words

- ERB opens on evidence — a real fingerprint match — not a hunch or a feeling that something looks wasteful.
- Five rounds, each with a different job: evidence, removal, and a Socratic checklist that ends in a named constraint or an honest admission.
- The question library is editable by the customer's own IE forever — wording and thresholds, no code required.
- The fingerprint's fields are the real ceiling — a genuinely new metric still needs a code change, and that's stated honestly, not hidden.

TALK IT THROUGH — 3 QUESTIONS FOR YOUR NEXT LINE WALK

1. Pick a step on your own line that's survived purely on “that's just how it's done” — could you actually name the real constraint if pushed to Round 5?
2. Find a repeated pattern (like our deflash example) across 3+ stations — has anyone ever traced it back to its real upstream cause?
3. If you were handing this tool to a customer's IE today, which trigger threshold would you expect them to want to edit first — and could they, without you?