

What Systemica Engineering is making

From: Joe Maxwell from Joe Maxwell
joesystemica@substack.com

To: joseph.maxwell02@proton.me
joseph.maxwell02@proton.me

On: Monday, 25 May 2026 at 0:58:37

Forwarded this email? [Subscribe here](#) for more

What Systemica Engineering is making

Pre-validation, model audit, and evidence packs before expensive action.

JOE MAXWELL
MAY 24



READ IN APP

What Systemica Engineering Is Building

Pre-validation, model audit, and evidence packs before expensive action

The last few posts have been more conceptual.

This one is more practical.

Systemica Engineering is the direction I'm building around a simple problem:

technical work often becomes expensive before people know whether the underlying assumptions are stable.

That can happen in engineering design.

It can happen in modelling.

It can happen in documentation.

It can happen when a project has enough information somewhere, but not in a form that lets people make a good decision.

So the practical aim is not to build a giant all-purpose system from day one.

The aim is smaller and more useful:

help turn messy technical work into structured evidence before people commit serious time, money, or energy to the wrong next step.

That means three main things.

Pre-validation.

Model audit.

Evidence packs.

The problem

A lot of technical work follows the same pattern.

Someone has an idea.

They build a model.

They run a simulation.

They make a design choice.

They write a report.

They move toward testing, fabrication, funding, patent work, or a bigger technical commitment.

But somewhere in that chain, weak assumptions can travel too far.

A model output looks clean, so it gets trusted.

A design option looks viable, but the governing constraint has not been found yet.

A technical document exists, but nobody can actually use it.

A feature looks important in a model, but nobody has checked whether it stays important when the model is stressed.

This is the gap Systemica Engineering is meant to sit in.

Not replacing expert judgement.

Not replacing formal validation.

Not replacing simulation or testing.

Sitting before those stages, where the work is still cheap enough to redirect.

TRS-GL: thermal pre-validation

The first practical runtime is called TRS-GL, or Thermal Runtime Studio.

The plain version:

TRS-GL helps screen early thermal design options before expensive simulation, prototyping, or testing.

It came from my individual engineering project on a furnace wall and removable viewing cartridge.

The design problem was simple enough to explain, but awkward enough to matter:

how do you keep a furnace wall insulation-dominated while adding a local viewing module that does not create a dangerous thermal bridge?

Instead of treating the first model as a final answer, TRS-GL treats the model as a controlled screening environment.

Declare the assumptions.

Run the thermal cases.

Sweep the parameters.

Find the governing case.

Classify options as viable, marginal, or non-compliant.

Export the evidence.

That is the point.

Not "the model is truth."

More like:

here is what was assumed,

here is what was tested,

here is which condition governed,

here is what failed,

here is what survived,

and here is what still needs formal validation.

That is pre-validation.

It does not replace Ansys, COMSOL, CFD, FEA, physical testing, certification, or engineering judgement.

It helps decide what is worth taking into those stages.

GIL/EDK: model audit

The second layer is GIL/EDK.

This sits more in the modelling and machine-learning world.

The plain version:

GIL/EDK checks whether model conclusions stay stable when modelling conditions are disturbed.

A model can perform well and still be fragile.

A feature can rank highly once and still collapse when the data split changes, noise is added, a control is introduced, or assumptions are perturbed.

That matters in scientific and engineering modelling because feature importance is often treated as explanation.

A descriptor looks important.

A graph looks convincing.

A ranking gets interpreted as a mechanism.

But static importance is only a snapshot.

The stronger question is:

does the feature keep behaving reliably under stress?

That is the shift from static feature importance to behavioural feature auditing.

Instead of only asking “what mattered in this model run?”, GIL/EDK asks:

what stayed stable?

what drifted?

what collapsed?

what only worked under one condition?

what survived negative controls?

what changed the follow-up decision?

That is useful because scientific models are not just judged by scores.

They are used to decide what to trust next.

Why evidence packs matter

The shared output across this work is the evidence pack.

An evidence pack is not just a result.

It is a structured record of how the result was produced.

For TRS-GL, that might include:

the design assumptions,

the material stack,

the thermal sweep,

the governing case,

the classification,

the limitations,
and the next validation recommendation.

For GIL/EDK, it might include:

the baseline model,
the perturbation trace,
the descriptor drift,
the variability,
the barcode state,
the negative-control behaviour,
and the claim ceiling.

The phrase “claim ceiling” matters.

It means the output should say what level of confidence it has earned.

Screening-level.

Prototype-level.

Audit-level.

Research signal.

Not validated.

Not certified.

Not ready for production.

That sounds boring, but it is where trust comes from.

A useful technical system should not just produce an answer.

It should tell you what kind of answer it is allowed to be.

The wider Systemica Engine

Behind these tools is a wider idea I call the Systemica Engine.

The simple version:

experiments should not be treated as isolated runs.

They should be treated as trajectories through a state space.

Most technical workflows still look like:

write code,

run experiment,

inspect result,

repeat.

That works, but a lot gets lost.

Why was this run done?

What assumptions changed?

What failed?

What recovered?

Which result was stable?

Which result looked good once but disappeared under stress?

The Systemica Engine is the long-term architecture for making modelling work more traceable, repeatable, and self-documenting.

But it is not the first thing to sell.

That is important.

The first useful tools have to stand alone.

TRS-GL should make sense as a thermal screening runtime.

GIL/EDK should make sense as a model audit layer.

Only later do they connect into a larger governed engineering runtime.

Start narrow.

Prove value.

Then connect.

The practical offer

So the practical work of Systemica Engineering is not "here is a giant system, please understand all of it."

It is closer to:

bring me a messy technical problem,

an early design,

a modelling result,

a scattered set of notes,

or a project that needs a clearer decision surface.

Then the work is to produce something more usable:

a structured technical note,

a pre-validation map,

a model audit,

a design-screening evidence pack,

a knowledge-base structure,

or a clear next-step recommendation with limits attached.

The aim is not to make decisions for people.

The aim is to make the decision surface clearer.

Why this matters

Expensive technical work should not begin from vibes.

It should not begin from a single clean-looking model output.

It should not begin from a document nobody can interrogate.

It should not begin from an assumption that has not been stressed.

Systemica Engineering is being built around the opposite habit:

before expensive action, make the system legible.

What is the claim?

What supports it?

What breaks it?

What governs the result?

What is still unknown?

What should be tested next?

That is the practical spine.

TRS-GL applies it to thermal design.

GIL/EDK applies it to model and descriptor stability.

The wider Systemica Engine applies it to governed experimentation.

Same pattern.

Different surfaces.

Better maps before heavier action.

SHARE

LIKE

COMMENT

RESTACK