

V1 Accuracy Is Becoming a Commodity The bottleneck in modelling has quietly changed.

From: Joseph.Maxwell
Joseph.Maxwell02@proton.me
To: Joseph.Maxwell02@proton.me
Joseph.Maxwell02@proton.me
On: Tuesday, 14 July 2026 at 18:25:51

Accuracy Is Becoming a Commodity

I. The bottleneck in modelling has quietly changed.

For much of the history of computational modelling, predictive accuracy has been the defining measure of progress. Better models were difficult to construct, expensive to evaluate, and often represented years of incremental advances in mathematics, computing and experimental understanding. Under those conditions it made perfect sense to judge progress primarily by predictive performance. If one model predicted reality more accurately than another, it was usually the more valuable scientific instrument. Improving prediction meant improving science.

That assumption is beginning to change.

Across engineering, materials science, biology, finance and increasingly every discipline touched by computation, we are entering an era in which producing another highly accurate predictive model is no longer the scarce resource it once was. Machine learning, automated model discovery, foundation models and rapidly expanding computational capability have fundamentally altered the economics of prediction. Tasks that once demanded months of specialist effort can now be explored in hours. Entire families of candidate models can be generated, optimised and compared automatically. In many domains, prediction is becoming increasingly abundant.

This is not because modelling has stopped improving. Quite the opposite. Modern models routinely discover statistical relationships that would have been inaccessible only a decade ago. They interpolate enormous datasets, uncover subtle interactions and increasingly contribute to scientific discovery itself. Their success, however, has quietly exposed a different limitation. As predictive performance becomes easier to obtain, its value as the sole measure of scientific progress inevitably begins to diminish.

Accuracy is becoming abundant.

What remains scarce is something else entirely.

Not prediction.

Not optimisation.

Not computation.

The ability to determine which predictions deserve our trust.

The title of this essay deliberately uses the word *commodity*, because commodities are not things that have lost their value. Steel is still valuable. Electricity is still valuable. Concrete remains indispensable to modern civilisation. What changes is not their usefulness but their scarcity. As production becomes cheaper and more routine, value migrates elsewhere in the system. I believe predictive modelling is beginning to undergo the same transition. The cost of generating another accurate model is falling rapidly. The cost of determining whether that model deserves engineering trust is not.

Every mature technology experiences a shift like this. Steam engines eventually stopped being constrained by steam itself and became constrained by metallurgy. Software development ceased being limited by our ability to write code and became limited by testing, version control and reproducibility. Manufacturing stopped asking whether parts could be produced and began asking whether they could be produced consistently. As production becomes easier, assurance becomes harder. The bottleneck moves from creating outputs to governing them.

Computational modelling appears to be reaching exactly the same point.

For decades, the central question was simple: *Can we build models that predict reality more accurately?* Increasingly, the answer is yes. Modern workflows can produce highly performant models with remarkable speed, often across problems that previously required years of manual effort. Yet that very success has created a new question, one that is arguably more difficult than the first.

Which of those increasingly accurate models actually deserve our trust?

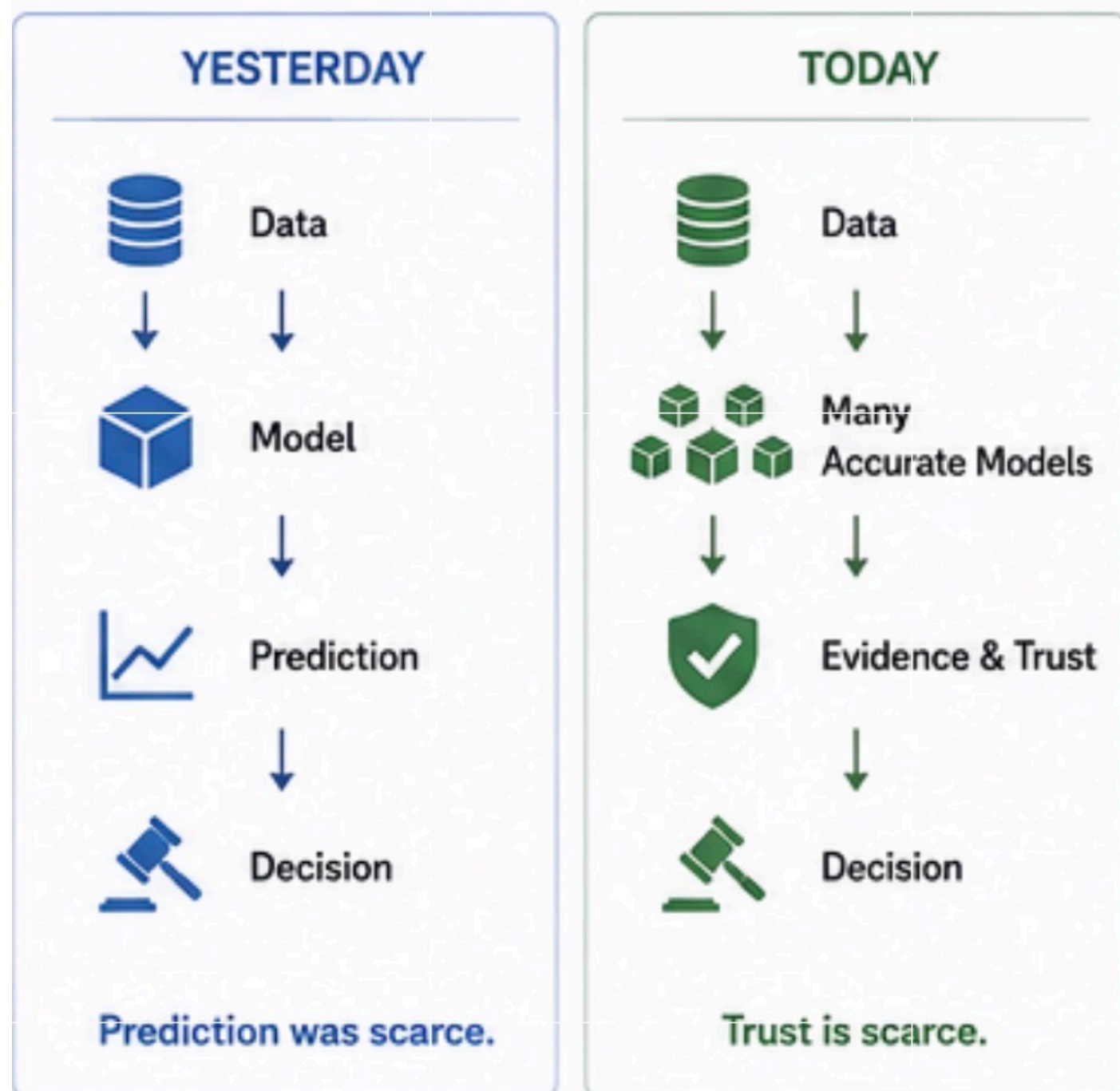
Prediction and trust are not the same thing. Prediction estimates what may happen under a particular set of assumptions. Trust determines what we are willing to believe, investigate, manufacture, publish or spend millions validating. Those decisions carry consequences far beyond a benchmark score. Research programmes are prioritised. Experiments are funded. Components are manufactured. Materials are selected. Clinical studies are advanced. Entire engineering programmes change direction because a model appears sufficiently convincing.

Yet the statistical performance of a model and the justification for acting upon it have never been the same quantity. We have simply become accustomed to treating them as though they were.

I suspect that the next decade of modelling will not be defined by another dramatic leap in predictive capability. Those improvements will undoubtedly continue, but they are no longer the only story. The more profound transition is quieter. Prediction is becoming increasingly automated. Judgement has not. The scarce resource is shifting from generating another answer to determining which answers deserve to become accepted knowledge.

That, I believe, is where the next bottleneck in modelling quietly waits.

Figure 1 — The Bottleneck Has Moved



Advances in computation and machine learning have shifted the limiting resource in modelling from generating predictions to determining which predictions deserve engineering trust.

Figure 1. The Bottleneck Has Moved

For decades the primary challenge in computational modelling was producing accurate predictions. As predictive capability becomes increasingly abundant, the limiting resource shifts towards determining which predictions deserve engineering trust.

I. Accuracy Was Never the Whole Story

If this argument sounds provocative, it shouldn't.

The claim is not that predictive accuracy has stopped mattering. It remains one of the most important properties any scientific model can possess. Without predictive capability there is nothing to validate, nothing to interpret and nothing to apply. Accuracy is the foundation upon which every other judgement rests.

The problem is that it has gradually become treated as though it were the entire building.

Two models may achieve almost identical predictive performance while possessing completely different scientific value. One may continue to reach the same conclusions after modest perturbations to the data, modelling assumptions or feature selection. The other may collapse as soon as those conditions change. Both may report the same R^2 . Both may satisfy the same benchmark. Yet only one represents a relationship that appears robust enough to support further engineering effort.

Imagine, for example, screening ten thousand candidate alloys for a lightweight aerospace component. Two independent modelling workflows identify the same material as a promising candidate. Both achieve similarly impressive predictive performance against historical data. At first glance they appear equally persuasive. Yet after modest perturbations—slight changes to the training data, alternative feature selections, different modelling assumptions or independent validation—one continues pointing towards the same alloy while the other begins recommending entirely different candidates. Both models remain accurate according to conventional benchmarks. Only one remains reliable enough to justify the next stage of expensive physical validation. The benchmark score did not reveal that distinction. The engineering workflow did.

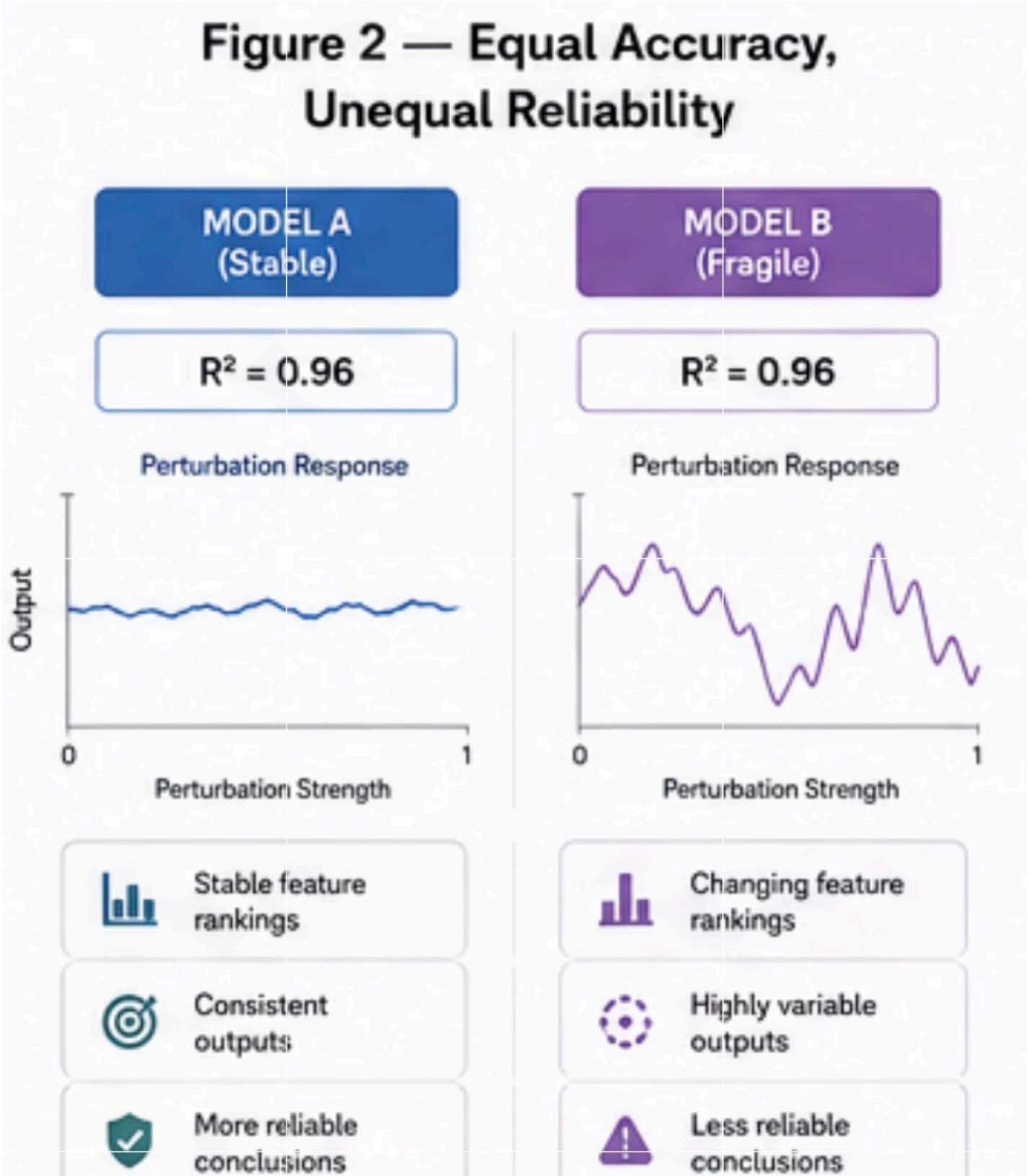
Traditional performance metrics struggle to distinguish between situations like these because they were never designed to answer that question. Metrics such as R^2 , RMSE, MAE or classification accuracy describe how well a model performed under a particular evaluation procedure. They tell us remarkably little about how confidently that result should influence engineering decisions beyond it.

This distinction has existed for decades. It has simply become more visible as prediction has become easier. When accurate models were rare, the practical objective was to obtain one. As accurate models become plentiful, comparison becomes inevitable. Once dozens of models can produce similarly impressive scores, the question naturally changes from *Which model predicts best?* to *Which prediction deserves to shape our understanding of reality?*

That transition is exactly what we would expect if predictive performance is gradually becoming a commodity. The benchmark remains essential. It simply ceases to be sufficient as the primary differentiator. Scientific maturity increasingly depends not on our ability to generate another prediction, but on our ability to determine what that prediction actually means.

That is not merely a statistical question.

It is an engineering question.



Predictive performance alone cannot distinguish between stable and fragile conclusions. Two models may achieve identical benchmark scores while exhibiting fundamentally different behaviour under perturbation.

Figure 2. Equal Accuracy. Different Reliability.

Predictive performance alone cannot distinguish between stable and fragile conclusions. Two models may achieve identical benchmark scores while exhibiting fundamentally different behaviour under perturbation.

Engineering has spent centuries learning that successful performance under ideal conditions is rarely sufficient to justify confidence. Modelling is beginning to rediscover the same lesson. The question is no longer whether we can predict reality. It is how we decide which predictions deserve to become part of it.

II. Engineering Already Knows the Answer

Engineering has encountered this problem before. In fact, it encounters it almost everywhere.

No responsible engineer certifies a bridge because a single finite-element simulation looked promising. Aircraft are not approved because one prototype completed a successful flight. New alloys are not introduced into critical applications because a tensile specimen happened to reach specification once. Every mature engineering discipline eventually learns the same lesson: reliable decisions require more than successful outcomes. They require evidence that those outcomes remain stable under variation, uncertainty and repeated challenge.

That philosophy quietly underpins almost every engineering assurance process. Fatigue testing exists because components rarely fail under ideal loading; they fail after years of accumulated stress. Process capability exists because manufacturing variation matters just as much as average performance. Safety factors exist because uncertainty cannot be eliminated, only managed. Redundancy exists because systems eventually encounter conditions that designers never explicitly anticipated. None of these practices replace good engineering. They exist because good engineering alone is rarely sufficient.

Scientific modelling has evolved many comparable techniques. Cross-validation asks whether performance survives different partitions of the data. External validation asks whether conclusions extend beyond the dataset that created them. Sensitivity analysis investigates how outputs respond to changing assumptions. Uncertainty quantification estimates confidence around predictions. Explainable AI attempts to expose the

relationships driving a model's behaviour. Robustness testing deliberately perturbs inputs to identify fragile behaviour. Bayesian approaches provide principled ways of representing uncertainty rather than hiding it behind point estimates.

Each of these methods represents an important step towards more reliable modelling. Collectively they have transformed computational science into a far more rigorous discipline than it was only a generation ago. None of them deserves to be dismissed; each solves an important part of the problem. If anything, their continued development is one of the reasons modelling has advanced so rapidly.

The point, however, is not that these techniques are missing.

The point is that they are usually treated as separate destinations rather than connected parts of a larger journey.

Cross-validation tells us one thing. Perturbation tells us another. Explainability offers a different perspective. Physical validation contributes something else again. Each produces valuable evidence, but those pieces of evidence are rarely organised into a coherent picture of what we should actually believe.

Modern modelling workflows have become remarkably good at collecting evidence. We gather benchmark scores, uncertainty estimates, feature rankings, sensitivity analyses, validation studies and increasingly sophisticated explanations of model behaviour. Yet those pieces often remain isolated from one another. We collect evidence without collecting confidence. We measure many things extremely well while still struggling to answer the simplest engineering question of all:

What should we actually believe?

That distinction becomes increasingly important as the number of available models continues to grow. If predictive accuracy is gradually becoming a commodity, then evidence itself becomes the differentiator. The challenge is no longer producing another successful model. It is understanding whether independent forms of evidence consistently point towards the same conclusion.

Imagine asking a room of engineers whether a bridge should be approved. One examines the stress analysis. Another reviews the fatigue calculations. A third inspects the manufacturing tolerances. A fourth studies corrosion behaviour. A fifth evaluates historical failures. None of them individually determines whether the bridge is safe. Confidence emerges because independent lines of evidence begin converging on the same conclusion. Engineering does not trust one measurement simply because it is impressive. It trusts the convergence of multiple independent observations.

Modelling is beginning to face an analogous challenge. As predictive systems become more capable, confidence itself becomes a multidisciplinary object. Reliability is no longer something that can be inferred from a single benchmark. It emerges from the interaction between multiple forms of evidence, each revealing a different aspect of how the model behaves when confronted with uncertainty.

Figure 3 — Reliability Is Accumulated



Reliable engineering decisions emerge from the accumulation of independent forms of evidence. Predictive accuracy remains essential, but it forms the foundation of confidence rather than its endpoint.

Figure 3. Reliability Is Accumulated.

Reliable engineering decisions emerge from the accumulation of independent forms of evidence. Predictive accuracy remains essential, but it forms the foundation of confidence rather than its endpoint.

In that sense, the future of modelling may look less like searching for the single best algorithm and more like building better maps of evidence. The objective shifts from asking *Which model performs best?* towards asking *What do independent sources of evidence consistently allow us to conclude?* That is a subtle change in language, but it represents a profound change in philosophy. We stop treating reliability as an attribute of the model itself and begin treating it as something that must be earned through convergence.

This distinction is subtle but important. A benchmark score is an observation. Confidence is an interpretation. They are related, but they are not interchangeable. Engineering has always understood this difference. A material passing one tensile test does not immediately become flight certified. A simulation matching one experiment does not automatically become accepted design truth. Confidence is earned because evidence continues pointing in the same direction despite repeated opportunities for it to fail.

Perhaps the most interesting consequence of this shift is that automation itself begins to change role. If generating predictive models becomes increasingly inexpensive, then producing another answer contributes relatively little value. The scarce resource becomes the ability to determine which answers deserve promotion into engineering decisions. In other words, automation no longer creates the primary bottleneck. It inherits it.

This is why I think discussions around artificial intelligence often become distracted by the wrong question. We ask how capable our systems are becoming. Increasingly, capability is not the limiting resource. Foundation models, automated machine learning and rapidly improving computational infrastructure continue pushing that frontier forward at remarkable speed. The more interesting question is whether our ability to organise, interpret and govern their outputs is improving at the same rate.

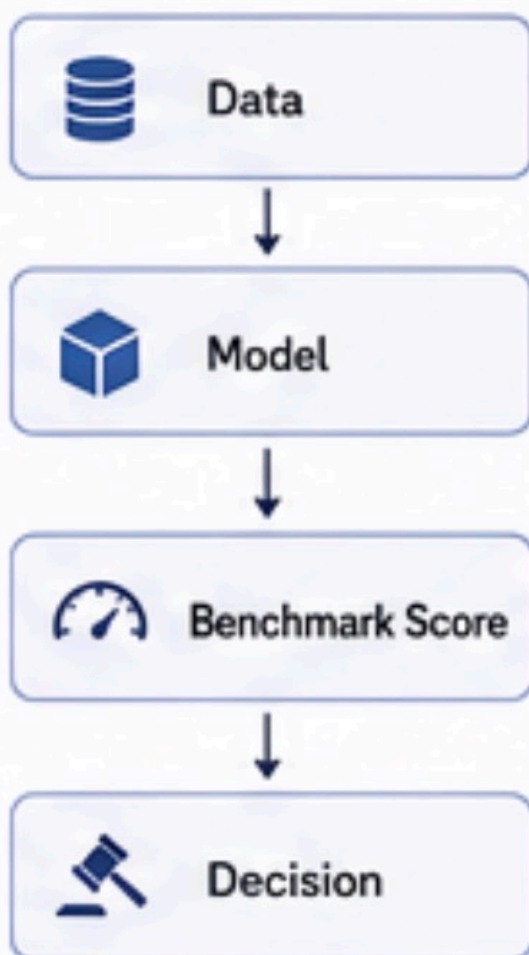
So far, I am not convinced that it is.

That observation has nothing to do with whether modern AI systems are impressive. They clearly are. It has everything to do with where value migrates as capability becomes abundant. Once prediction becomes relatively inexpensive, engineering effort naturally shifts towards understanding, validating and governing those predictions. The bottleneck quietly moves again.

That, I suspect, is where modelling is heading. We are moving away from an era in which the central challenge was producing accurate outputs and towards one in which the central challenge is governing those outputs responsibly. Prediction becomes one stage in a much larger engineering workflow. The real objective is no longer simply to minimise error. It is to maximise justified confidence.

Figure 4 — From Model Evaluation to Evidence Mapping

CONVENTIONAL WORKFLOW



EMERGING WORKFLOW



As prediction becomes increasingly abundant, the bottleneck shifts towards organising multiple forms of evidence into a coherent basis for decision-making.

Figure 4. From Model Evaluation to Evidence Mapping.

As prediction becomes increasingly abundant, the bottleneck shifts towards organising multiple forms of evidence into a coherent basis for decision-making.

If this is where modelling is heading, then the question becomes less about inventing another predictive algorithm and more about designing the infrastructure that allows evidence itself to move through a modelling workflow. Over the last few years, that question gradually became the centre of my own work.

III. Mapping Reliability

Over the last few years I found myself becoming less interested in building another predictive model and increasingly interested in understanding how evidence itself should move through a modelling workflow. Somewhere along the way I realised that I had quietly stopped asking *“Can I improve this prediction?”* and started asking something that felt much more fundamental:

“What would it actually take for me to trust this result?”

At first that sounds like a philosophical question.

In practice it turns out to be a deeply practical engineering problem.

Trust is rarely created by a single observation. It accumulates. Every additional experiment, every successful perturbation, every independent reproduction, every physically plausible explanation and every modelling approach that converges on the same conclusion contributes something small to a much larger picture. Conversely, every unexplained sensitivity, every contradictory result and every hidden dependency should reduce our confidence accordingly. Reliability behaves less like a binary property and more like a landscape that gradually becomes visible as different forms of evidence accumulate.

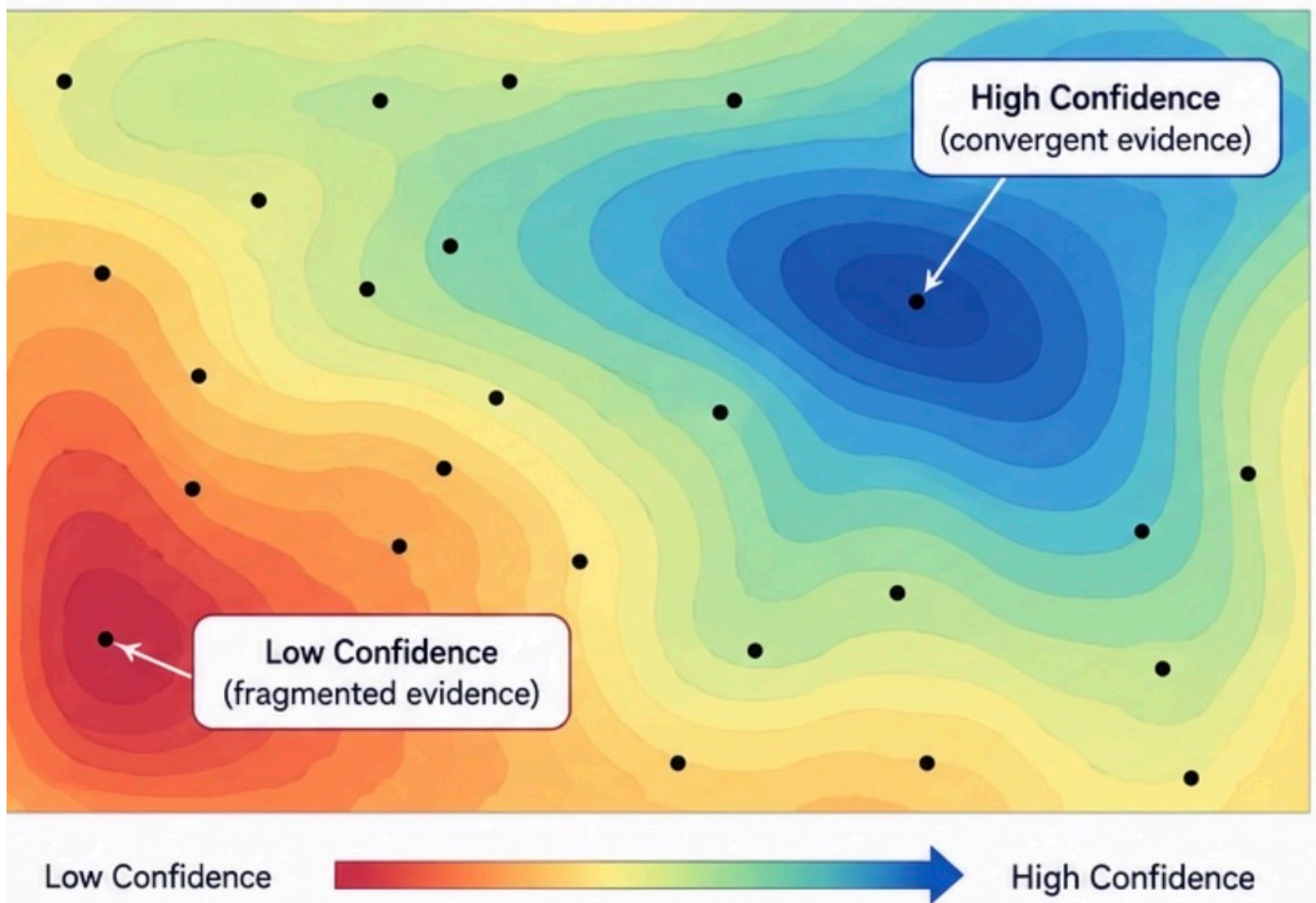
That observation changed how I viewed many of the techniques I had been using for years. Perturbation was no longer simply a robustness test; it became a way of exposing how stable a conclusion really was. Cross-validation was no longer just another performance check; it became one independent source of evidence among many. Explainability, sensitivity analysis, physical constraints, reproducibility and external validation all began to look less like separate methodologies and more like different perspectives on the same underlying question.

None of them individually determined what should be believed.

Together they started forming something much more useful.

A map.

Figure 5 — Mapping Reliability



Reliability is represented as a landscape rather than a single score. Different forms of evidence gradually reveal regions where conclusions remain stable and regions where they become increasingly uncertain.

Figure 5. Mapping Reliability.

Reliability is represented as a landscape rather than a single score. Different forms of evidence gradually reveal regions where conclusions remain stable and regions where they become increasingly uncertain.

That realisation slowly changed the objective of the work itself. Instead of asking how another model could be made slightly more accurate, I became increasingly interested in how independent forms of evidence could be organised into something that was navigable. Not simply collected, but structured. Not merely stored, but connected. Not just available, but interpretable.

The more I explored that idea, the more obvious it became that modern modelling possesses remarkably sophisticated tools for generating evidence, but comparatively little infrastructure for organising it. We have libraries that train models, explain models,

validate models, visualise models and deploy models. What we possess far less frequently is an engineering framework that asks a much simpler question:

Given everything we currently know, what conclusion is actually justified?

That observation gradually became the foundation of what is now the **Systemica Atlas**.

Despite the name, the Atlas is not intended to compete with existing modelling frameworks, optimisation libraries or machine learning platforms. Those tools already perform their respective jobs exceptionally well. The Atlas attempts to address a different layer of the problem entirely. It is an attempt to organise evidence itself.

Rather than asking *“Which model performs best?”* it asks *“What can we currently justify believing, and why?”*

Rather than treating validation as something that happens after modelling, it treats validation as something that accompanies every stage of the modelling process.

Predictions are generated.

They are challenged.

They are perturbed.

They are compared.

They are reproduced.

They are examined against physical understanding.

Only then are they promoted and only to the extent that the accumulated evidence justifies doing so.

Figure 6 — From Prediction to Evidence



The Atlas views modelling as the progressive accumulation of evidence rather than the production of isolated predictions. Confidence is earned by surviving successive forms of challenge.

Figure 6. From Prediction to Evidence.

The Atlas views modelling as the progressive accumulation of evidence rather than the production of isolated predictions. Confidence is earned by surviving successive forms of challenge.

This distinction may appear subtle, but I believe it changes the role of automation quite profoundly.

Much of today's conversation around artificial intelligence focuses on making systems more capable. My own work gradually moved in almost the opposite direction. I became increasingly interested in how capable systems remain **honest**.

How do we ensure uncertainty travels alongside a prediction instead of disappearing the moment a benchmark score is reported?

How do we preserve provenance as workflows become increasingly automated?

How do we distinguish between outputs that merely exist and outputs that deserve to influence engineering decisions?

How do we ensure that evidence never becomes detached from the assumptions that produced it?

Those questions are less about intelligence than governance.

They concern how evidence moves through increasingly automated technical systems without becoming detached from the reasoning that justifies it.

Eventually I realised that the work had gradually converged on a surprisingly simple objective.

Making honest automation efficient.

Not making automation more intelligent.

Not making automation more autonomous.

Making it more honest.

Because honesty, in engineering, is simply the discipline of ensuring that conclusions never become detached from the evidence that justifies them.

Efficiency matters because none of those safeguards are useful if they make technical work prohibitively expensive. Scientists should not have to choose between moving quickly and being rigorous. Engineers should not have to sacrifice provenance in order to explore design spaces. Researchers should not have to rebuild evidence every time a workflow changes.

The objective is not to replace scientists, engineers or modellers.

It is to reduce the cost of doing careful work.

That idea has gradually grown beyond modelling itself. The same pattern appears in software engineering, repository management, design exploration, technical governance and AI-assisted development. Modern automation is remarkably good at producing outputs. What it still struggles to do is communicate the reliability of those outputs with the same efficiency.

We have spent decades making generation cheaper.

We are only beginning to make justified confidence scalable.

Whether the particular ideas behind the Atlas ultimately prove useful is, in many ways,

secondary to the broader transition that motivated them.

I suspect modelling itself is entering a period in which **evidence becomes the primary engineering resource**.

Predictive capability will continue improving.

Models will become larger.

Training will become cheaper.

Automation will become increasingly routine.

That trajectory appears almost inevitable.

The harder challenge will be ensuring that our ability to justify conclusions grows at the same pace as our ability to generate them.

Every mature engineering discipline eventually discovers that producing results is only half the problem.

The other half is determining which results deserve to become accepted knowledge.

Computational modelling is quietly arriving at that moment now.

The future of modelling will not be won by whoever produces the most predictions.

It will be won by whoever most efficiently determines **which predictions deserve to become knowledge**.

That is the bottleneck I believe has quietly emerged.

Accuracy is becoming a commodity.

Reliable judgement is not.

Figure 7 — The Programme Philosophy

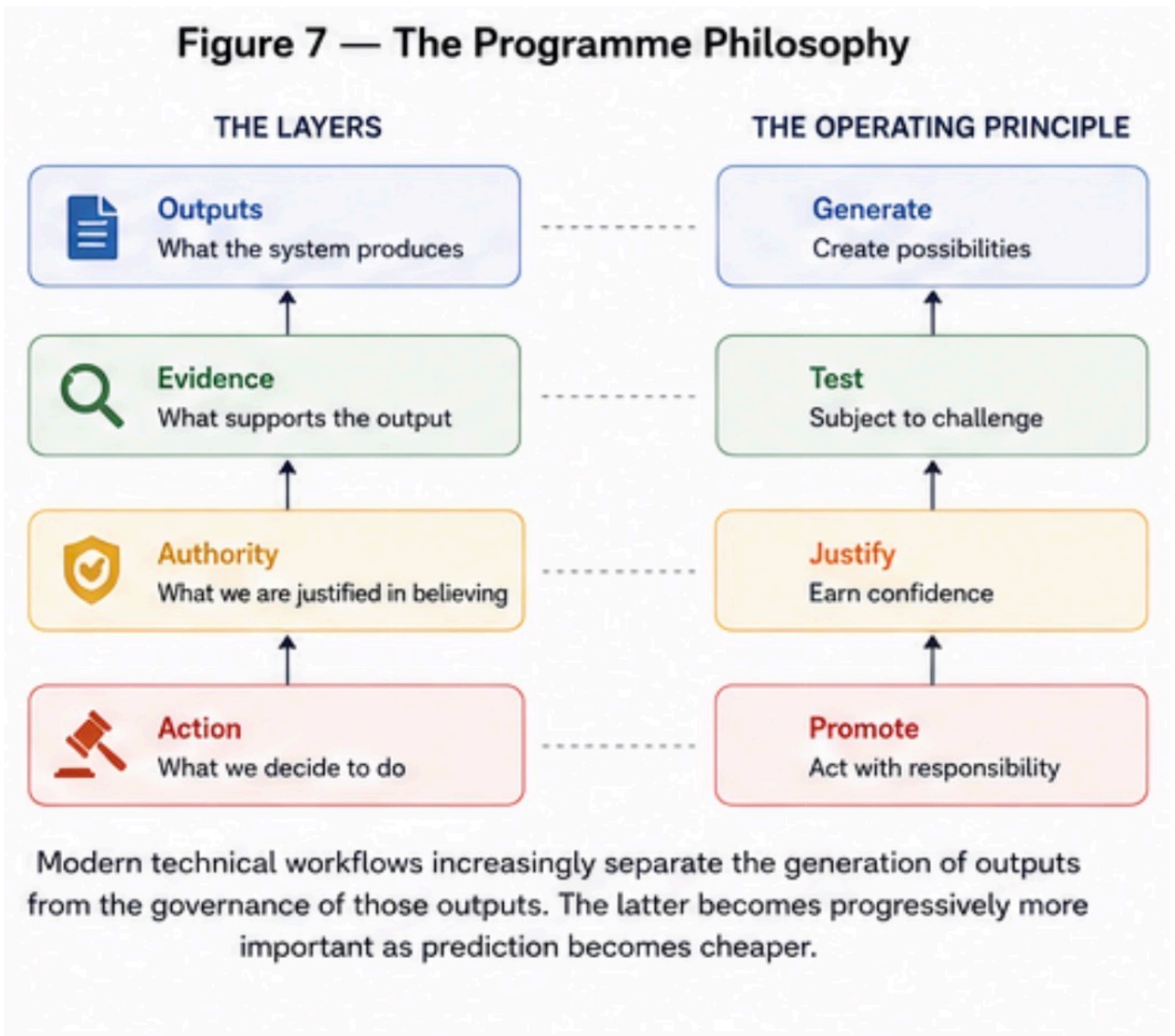


Figure 7. Governing Outputs.

Modern technical workflows increasingly separate the generation of outputs from the governance of those outputs. As predictive capability becomes increasingly abundant, engineering value migrates towards organising, validating and communicating evidence rather than simply producing another prediction.

End Note

Prediction will continue improving.

Models will become larger.

Automation will become cheaper.

Entire scientific workflows will become increasingly autonomous.

None of that removes the need for judgement.

If anything, it makes judgement more valuable than ever.

Every mature engineering discipline eventually reaches the point where producing results becomes easier than deciding which results deserve trust. Computational modelling is quietly arriving at that moment now.

As my mum has always said,

"Honesty is the best policy."

Perhaps it's finally time we took that a little more literally.