

Three Kinds of Wrong

The Structural Shape of AFS's Exposure to Error

By Mario Aiello – Adaptive Ways

Most approaches worry about being wrong. They worry about it the same way. Am I right or am I wrong. Did I get it or did I miss it. Like “wrong” is one thing.

It isn't.

Dave Snowden makes the point in passing. There are three kinds of wrong a model can be. Bounded but coherent. Incoherent. Or coherent and misapplied. Most people nod and move on. The nod is the problem. The three kinds aren't just three labels. They're three different shapes of trouble, and the defences you build against one of them tend to make another one worse.

This paper is about what that means for the AFS.

The short version: AFS is well-positioned in the first kind, structurally exposed to the second, and operationally exposed to the third, and the moves it makes to defend against the third are the same moves that increase its exposure to the second. There's no clean fix. There's only navigation.

That sounds like bad news. It isn't. It's just specificity. Most approaches, when they fail, fail confused about why. AFS can do better than that. Not because it's more right. Because it's more honest about the particular shape of how it might be wrong.

The three kinds

Snowden's distinction is small enough to write on a napkin and load-bearing enough to organize a research program around. So it's worth getting clean before applying it to anything.

Kind one: bounded but coherent. A model can be imprecise, simplified, or limited in scope, and still be internally coherent, consistent with the evidence inside its domain, and properly bounded. Newtonian mechanics is the standard example. It's wrong, relativity replaces it, but it's wrong in a useful way. It works for the speeds and scales most people will ever encounter. The mistake isn't in the model. The mistake would be in mistaking the model for the territory.

This is the most respectable kind of wrong. It might be the only kind a working approach can plausibly aim for. Every map simplifies. The question is whether the simplification is honest about itself.

Kind two: incoherent. A second kind of wrong is incoherence. Internal contradictions. Category errors, talking about a process as if it were a state, or a capacity as if it were a behavior. Confusions between levels of description, where a claim about individuals slides into a claim about systems without anyone noticing the slide. Or claims that can't be tested even in principle, because every observation can be reabsorbed into the model.

Incoherent models can be internally pleasing. They can feel comprehensive. That's part of the problem. The pleasure of completeness can mask the absence of structure. A model that explains everything explains nothing, and a vocabulary that means whatever the speaker needs it to mean isn't a vocabulary, it's a mood.

Kind three: coherent but misapplied. The third kind is the coherent model used outside its valid domain. The model itself is fine. The application isn't. This is what happens when something that worked in one context gets exported to another context where the assumptions don't hold, and nobody notices because the vocabulary still travels.

Lean, applied to a context that isn't waste-driven. Agile, applied to a context that isn't iteration-friendly. A diagnostic instrument calibrated for one kind of organization, used on another. The model doesn't fail because it's broken. It fails because it's elsewhere.

These three kinds aren't a hierarchy. They're a topology. Different exposures, different defenses, different consequences. With them in hand, we can look at what the AFS is doing and what the AFS isn't doing.

AFS through each kind

Where AFS is designed to live

Kind one is the kind AFS is built for. The whole architecture announces it. The Four Questions for Fitness are not a complete theory of organizations; they're four questions. The Adaptive Signal Field doesn't claim to capture every dimension of organizational life, or orient with certainty; it senses and surfaces without prescribing. The Friction-Fitness Map validates without certifying.

The approach/framework distinction, made explicit in The Canon Problem, is the strongest single signal that AFS knows it's operating in kind one. A framework presents itself as portable, repeatable, complete. An approach declares that it requires practitioner judgment, that it doesn't fit every context, that the user is part of the instrument. Those declarations are kind-one declarations. They are the model saying, in advance, where it stops.

Inside its scope, AFS is reasonably coherent. The vocabulary holds together. Fit for purpose, fit for context, fit for execution, fit for improvement, these aren't redundant and they aren't overlapping in ways that collapse the distinctions. The diagnostic instruments produce readings

that practitioners can act on. The Stress-Test paper exists, which means the system has at least once tried to falsify itself in writing.

This isn't the interesting part of the analysis. Kind one is where AFS is comfortable. The interesting parts are kinds two and three.

The incoherence exposures

Kind two is the kind AFS has to actively police, because approaches that rely on practitioner judgment can drift into incoherence without anyone noticing. The drift doesn't announce itself. It accumulates.

Three exposures are worth naming.

The first is category slippage. The word fitness does at least three jobs in AFS. Sometimes it names a state — this organization is fit for its context. Sometimes it names a process, this organization is becoming fit, or losing fitness over time. Sometimes it names a capacity, this organization has the means to stay fit through whatever comes next. These aren't synonyms. A state can be measured at a moment. A process unfolds over time. A capacity is a disposition that may or may not be exercised. F-DD does some work to bind these together as a single design logic, and the binding is plausible. It's also worth checking, periodically, that the binding still holds, that fitness in one part of the canon means the same thing as fitness in another part, and that the slide between meanings is being made by argument rather than by drift.

The second is level confusion. AFS's Portfolio Management (AFS-PM) operates at the strategic level, orchestrating competing organizational WHYs. Simple Agility operates at the execution level, where work meets reality. The Adaptive Signal Field operates at a meta-diagnostic level, helping practitioners locate themselves. These are three different scales. When a claim about portfolio-level orchestration travels to the execution level without explicit translation, or when a meta-diagnostic observation gets cashed out as an execution-level prescription, the system can look like it's saying something coherent that turns out, on inspection, to be a slide between scales. The defence is to keep the layers visible. The risk is that, in practice, the layers blur, especially when a single practitioner is moving between them in a single engagement.

The third is the testability problem, and this one is the hardest. What would count as AFS being wrong? The Friction-Fitness Map validates within AFS. Reverse Validation tests against AFS's own logic. The Stress-Test paper applies AFS-style critique to AFS itself. These are honourable moves, and they do work that purely external critique can't do, they reach the parts of the approach that only the approach can see. But they are also recursive. The system is partly validating itself. If every counter-evidence can be reabsorbed as another AFS reading, the practitioner misapplied the instruments, the context wasn't read correctly, the diagnostic

conditions weren't met, then the approach is doing the thing kind two warns about. It's becoming a vocabulary that means whatever it needs to mean.

This is the price of practitioner-judgment-as-feature. It buys flexibility. It buys context-sensitivity. It buys the explicit refusal to be a framework. What it costs is sharp falsifiability. The cost is real, and pretending it isn't real is itself a category error.

The misapplication risk

Kind three is the operational risk. AFS was developed in a particular kind of organization, knowledge work, mostly post-industrial settings, mostly European financial services and adjacent contexts. The diagnostic instruments assume a certain kind of organization: one that can engage in conversation about its own fitness, one with the slack to host that conversation, one whose work product is shaped enough by knowledge-work patterns that the AFS vocabulary recognizes it.

Outside that envelope, the application becomes a question.

A manufacturing line is a different beast. The fitness conversation that AFS can host in a banking transformation may not survive contact with a context where the work is physical, the cycle times are short, and the variability is in materials rather than meaning. An emergency response organization operates in a register where deliberation is itself a failure mode. A very early-stage startup may not have enough organization yet to be diagnosed at all. None of this means AFS is wrong in those contexts. It means AFS is somewhere else.

The borrowed components carry their own envelopes too. Lean has a domain. Theory of Constraints has a domain. Real Options has a domain. Flow has a domain. AFS inherits those domains whether it acknowledges them or not. When Lean travels into a context where waste isn't the binding constraint, Lean misapplies. When AFS imports Lean and travels somewhere Lean shouldn't go, AFS imports Lean's misapplication too. The composability of the toolkit is one of its strengths. It is also a vector for kind three.

The defence AFS has built against this is the practitioner-judgment stance. It's an approach, not a framework. The practitioner reads the context and adapts. The vocabulary is a starting point, not a prescription. This works. It's the right defence.

It's also what brings us to the trade-off.

The trade-off

Here is the thing the rest of the paper exists to deliver.

AFS's primary defence against kind three, the misapplication risk, is the practitioner-judgment stance. The approach travels well because it travels light. It doesn't import its assumptions as

rules; it imports them as questions. The practitioner is the instrument that adapts the instrument to the context. When the context shifts, the instrument shifts with it.

That defence works. It is, in fact, the right defence. A more rule-like AFS, one that shipped with mandatory diagnostics, with required sequences, with conditions for use written out in protocol form, would travel into more contexts unmodified and would misapply in more of them. The practitioner-judgment stance is what keeps AFS from becoming the thing it warns against.

But the same defence increases exposure to kind two.

If every application is bespoke practitioner judgment, falsifiability gets thin. If the vocabulary's meaning is always tuned to the situation, category slippage becomes harder to police, because there's no canonical category to slip from. If the validation instruments are themselves part of the approach being validated, the loop closes. The very flexibility that protects against kind three is what creates the conditions under which kind two can drift in unnoticed.

Now run the move in the other direction. Tighten the conceptual machinery. Define fitness more precisely and lock the definition. Make the Adaptive Signal Field a checklist with scoring criteria. Specify the sequence in which the Adaptive Signal Field, the Friction-Fitness Map, and the AFS Adaptive Engine must be applied. Publish the conditions under which each instrument is and isn't valid. Do all of this, and the kind two exposure shrinks substantially. The category slippage gets policed by the definitions. The level confusion gets policed by the protocols. The testability problem gets traction, because now you can specify what would count as AFS not working.

And the kind three exposure expands. Because rule-like portability is exactly what travels into domains where the rules don't hold. The more AFS becomes a framework, repeatable, codified, complete, the more confidently it gets carried into contexts where its assumptions are quietly wrong, and the more those contexts get read through a vocabulary that doesn't fit them. The defences against kind three depend on the practitioner being able to say not here, not like this, and the more the approach hardens into procedure, the harder that judgment becomes to exercise.

This is not a flaw to fix. It's a structural fact. You can't close both flanks at once. Tightening one loosens the other.

What this means, practically, is that the choice between kind two and kind three exposure isn't a choice that gets made once. It gets made continuously, across publications, in design decisions about new instruments, in how engagements are run, in how the approach is taught. Every time AFS hardens, adds a definition, formalizes a sequence, codifies a step, it trades some kind three exposure for some kind two protection. Every time AFS softens, adds a caveat, restores practitioner discretion, refuses to specify, it trades the other way.

The honest move is not to pretend either defence is sufficient on its own. The honest move is to navigate the trade-off out loud. Make the trade visible. Let practitioners and readers see, in any given component, which side of the trade it's sitting on, and why.

This is also why the practitioner-judgment stance is not a license. If practitioner judgment is the defence against kind three, then practitioner judgment carries a corresponding burden, the burden of policing kind two without the help of rules. The flexibility comes with maintenance.

Limitations

A few things this analysis doesn't do, and shouldn't be read as doing.

Snowden's three-kinds distinction is itself a model. It's a useful one. It's not the only one, and it isn't comprehensive. There are kinds of wrong it doesn't name. *Sociological wrongness*, when a model is right but no one will use it, or when it's adopted for the wrong reasons and produces the wrong outcomes despite being internally fine. *Temporal wrongness*, when a model was right and isn't anymore, because the context has moved while the vocabulary has stayed. *Ethical wrongness*, when a model works, in the technical sense, but shouldn't, because what it makes possible isn't something that should be possible. None of these are addressed here. They are real, and a complete audit would name them.

The recursion problem doesn't go away just because the lens is external. Using Snowden to audit AFS doesn't escape the validation problem. It relocates it. Snowden's taxonomy itself sits inside a larger set of assumptions about models, knowledge, and applicability. Those assumptions can be argued with. Anyone who finds Snowden's three-kinds insufficient as a way of thinking about model error will find this paper insufficient as an audit. Fair enough.

This paper is also written from inside AFS. It is not a neutral assessment by an outside observer. It is a self-examination, conducted by a practitioner of the approach, using a lens borrowed from elsewhere. The practitioner-as-auditor has access that an outsider doesn't, and blind spots that an outsider wouldn't have. Both should be assumed.

Finally, the trade-off described here is a tendency, not a theorem. The claim is that defences against kind three tend to increase exposure to kind two, and vice versa. There may be specific design moves that improve both flanks at once, or that don't trade as steeply as the general pattern would suggest. Those moves are worth looking for. The trade-off isn't a wall. It's a slope that has to be climbed deliberately.

Bridge

What changes after this?

Not the components. The Adaptive Signal Field is still the Adaptive Signal Field. The Friction-Fitness Map, the AFS Adaptive Engine, F-DD, the four Level 2 capability domains, none of them are different on the other side of this paper than they were before it.

What changes is how AFS gets talked about when something goes wrong.

When an engagement underdelivers, when a diagnosis turns out to have been mistaken, when a recommendation fails to land, the first question becomes which kind of wrong. Was the model bounded and used inside its bounds, just imprecise about a particular detail? Was it incoherent, a category slip, a level confusion, a piece of vocabulary doing too many jobs at once? Or was it coherent and elsewhere, applied to a context where the assumptions didn't hold, regardless of whether anyone saw that at the time?

These three questions don't excuse failure. They locate it. And locating failure is the precondition for learning from it without either abandoning the approach prematurely or defending it past the point of plausibility.

Over time, if the three-kinds vocabulary holds, it can become part of the practitioner's standard equipment. The way Cynefin domains became shorthand for talking about decision contexts, three kinds of wrong can become shorthand for talking about model exposure. That sounds like a kind two problem. That's a kind three risk we should probably name. That move trades kind three exposure for kind two protection, are we sure that's the trade we want here? Vocabulary at this level is not decoration. It changes what gets noticed.

Most approaches, when they fail, fail confused about why. The confusion isn't an accident. It's what happens when a single category, wrong, is asked to do the work of three. Splitting the category doesn't make AFS more right. It makes AFS more legible to itself when it's wrong.

Which is a different kind of progress. And possibly a more durable one.

Part of the AFS Open for Review series — adaptiveways.site/afs/open-for-review