

What It Takes to Lift the Enduring Constraint

BY JAY DUNNE · SERIES Part 3 of 3 · DOMAIN Enterprise Risk Intelligence

This is the third and final installment in a series on why enterprise risk intelligence keeps underperforming the investment made in it. Part 1 diagnosed the pathology firsthand: a C-suite disconnect, risk domains that don't talk to each other, and a function perceived as a cost center rather than a strategic asset. Part 2 pushed past those symptoms to name the actual constraint — not analytic capability, but organizational architecture: the decision rights, the reporting lines, the executive sponsorship, and the cross-domain governance that determine whether analysis produced anywhere in the enterprise ever reaches a decision. That architecture is what no procurement decision can buy.

Part 2 of this series closed on that note. AI makes the constraint impossible to keep hiding behind an analyst-hour excuse — the capability is arriving faster than most organizations' readiness to use it. That's the gap Part 1 already implied without naming directly: the question the field has been avoiding is not whether the capability will arrive, but what the organization has to become to use it.

This piece answers that question directly. Not with a product, and not yet with an architecture — that comes next. Here, the task is narrower and more disciplined: naming what the new function must be capable of, and what the organization around it must provide, before any specific structure can be evaluated against those requirements.

The instinctive fix to a silo problem is to add coordination. More meetings between domain leads. A shared dashboard. A quarterly cross-functional review. Enterprises have been trying versions of this for a decade, and the pathologies Part 1 documented — the C-suite disconnect, the domain silos, the cost-center perception — have not moved. That persistence is itself the evidence that coordination is the wrong fix for the problem being solved.

Organizational scholarship on cross-boundary knowledge work gives a precise account of why. Paul Carlile's research on knowledge boundaries in product development distinguishes three kinds of difficulty that can separate specialized groups. The simplest is syntactic — different vocabularies, resolved by agreeing on shared terms. More difficult is semantic — shared vocabulary that still carries different meanings, resolved by translation and shared interpretation. Hardest is pragmatic — differences that are not just about language or meaning but about what each group has invested in its own way of solving problems, where accommodating another group's knowledge means changing your own at real cost.

Cyber security teams have often invested years building threat frameworks that resist reinterpretation by outsiders. Geopolitical analysts have invested in country risk models that don't translate cleanly into financial terms. Financial risk functions have invested in quantitative approaches that treat qualitative geopolitical input as a footnote — acknowledged in the risk register, a line in a larger report, rarely something that changes the underlying number. Corporate security and executive protection functions have invested in physical security and duty-of-care frameworks whose reporting language reads to executive leadership as tactical rather than strategic — a mismatch that has nothing to do with the quality of the work and everything to do with translation. These are not communication failures. They are pragmatic boundaries, in Carlile's sense — and pragmatic boundaries are not resolved by adding a meeting. They require transforming the knowledge itself, which is a fundamentally different kind of work than transmitting it.

This is the first thing the new function has to be: structurally distinct from the domains it integrates, not an overlay coordinating between them. A coordination layer assumes the domains simply need to talk more. A synthesis function assumes talking is not the bottleneck — transformation is, and transformation requires a function built for that purpose, with its own mandate, rather than borrowed time from people whose actual job is defending their own domain's framework.

Naming the function's role does not yet tell you what it needs to be capable of, because not every cross-domain risk problem is the hardest kind. Some genuinely are syntactic — a shared taxonomy would resolve them. Some are semantic — shared interpretation would lead to resolution. Treating every integration failure as if it requires the full weight of pragmatic transformation is its own kind of category error, and it's a costly one: it either over-invests in transformation work where a simpler fix would do, or — more commonly in practice — it under-invests, because the organization defaults to the cheapest available fix (a shared dashboard, a taxonomy) for problems that actually require the harder kind of work.

Carlile's later work formalizes this distinction. Boundaries can be gauged along two properties: difference, the distance between how two groups understand and solve problems, and dependence, how much each group's work constrains the other's. Where both are low, a shared syntax is often sufficient — this is *transfer*. As novelty and interdependence increase, translation becomes necessary — *translate*. At the boundary's most complex, differences and dependencies generate consequences that can only be resolved by changing the knowledge itself — *transform*.

This gives the new function its first strategic requirement, and it is a diagnostic one: the capacity to correctly classify which kind of boundary a given cross-domain risk problem actually presents, before applying a fixed methodology to it. This sharpens rather than restates Part 2's thesis. Architecture is the binding constraint — but architecture has to be differentiated, capable of applying transfer, translation, or transformation as the problem actually demands, not a single integration layer applied uniformly regardless of what's underneath it.

Given that diagnosis, four specific capabilities follow. These are not aspirational, but rather drawn directly from what Carlile's research identifies as the requirements of an effective boundary process, mapped onto what a synthesis function in enterprise risk actually needs to do.

Shared syntax capacity. The function needs the ability to establish common reference points across domains — shared terms, shared units of measurement for risk exposure, shared

formats for representing findings — without pretending the domains are the same discipline wearing different vocabulary. This is foundational but insufficient on its own; it is necessary in every case and sufficient only in the simplest ones.

Capacity to specify differences and dependencies. Beyond shared language, the function needs the ability to surface where domains actually depend on each other's outputs — where a cyber finding changes the financial exposure calculation, where a geopolitical shift changes the assumptions underneath a supply chain risk model. This is translation work: making explicit what each domain would otherwise leave implicit, because implicit assumptions are exactly what fail silently when domains operate in isolation.

Capacity to transform domain knowledge into decision-relevant synthesis. This is the pragmatic-boundary payload, and it is the capability current risk functions are most structurally missing. It is not enough to represent each domain's findings side by side for an executive to reconcile — that is what most cross-functional dashboards already do, and it is precisely why they haven't solved the problem. Side-by-side representation is a transfer-level fix applied to a transformation-level boundary. The function has to do the reconciling itself — which means it has to be capable of altering how a domain's own finding gets framed, not just relaying it. A cyber team's technical severity rating has to become a financial exposure figure. A geopolitical analyst's probability assessment has to become a supply chain contingency threshold. Neither of those translations is native to the domain that produced the original finding, and neither can be outsourced back to the executive receiving the synthesis — that defeats the purpose of having a synthesis function at all. This is the capability that separates a function actually doing integration from one simply aggregating reports under a shared cover page.

Iterative capacity. The requirement most likely to be skipped, and the one with the sharpest evidence behind it. Carlile's own case work is instructive here: in a vehicle redesign program, two platforms used the identical modeling tool and converged on the identical numeric target for a key design parameter — and one produced a smooth downstream launch while the other generated significant costly rework. The difference was not the tool. It was that one team let the boundary process run its full iterative course, and the other locked in a solution too quickly, having skipped the harder work of representing and negotiating differences before converging. The lesson transfers directly: deploying the same AI-enabled synthesis capability across two enterprises will not produce the same outcome if only one of them builds standing iteration

into how the function operates, rather than treating integration as a one-time project with a defined end date.

The four capabilities above describe what the function must be analytically capable of. They say nothing about whether the institution around it can actually sustain that capability under real conditions — not steady-state monitoring, but the moment a genuine cross-domain crisis actually hits.

Research on cooperation during transboundary crises is instructive here, even though its native context is intergovernmental rather than corporate. Arjen Boin is among the most-cited scholars in crisis management research, and his work with Donald Blondin extends that scholarship to a specific question: why transboundary cooperation happens in some crises and not others. Blondin and Boin's framework identifies three institutional capacities that determine whether coordinated response actually happens when it's needed, rather than remaining an aspiration. The first is analytical capacity — accurate, authoritative synthesis that lets different actors agree on the cause, scope, and urgency of a threat. Their finding is pointed: this capacity often already exists somewhere in the system. It just sits scattered across institutions and policy sectors that were never built to combine it. That is close to an exact restatement of Part 2's central claim, arrived at independently, in a different field entirely.

The second capacity is formal agreement, established before a crisis, not improvised during one — pre-committed authority rather than authority negotiated under pressure. The closest real-world analogues are treaty commitments made for exactly this reason: NATO's Article 5 and the EU's Solidarity Clause both exist to settle, in advance, who acts and under what authority, precisely so that question doesn't have to be litigated in real time while a crisis is unfolding. The third capacity is standing coordination infrastructure — mechanisms built to facilitate collaboration in fast-moving conditions, rather than an ad hoc arrangement assembled after something has already gone wrong.

Translated into enterprise terms, this is the new function's second set of requirements, this time institutional rather than analytical: synthesis capacity that is not scattered across risk domains but actually combined; executive sponsorship and decision authority established in advance of a crisis, not improvised in the middle of one; and standing infrastructure for cross-

domain collaboration, rather than a project-based integration effort stood up reactively once the cost of not having it has already been paid. A function can have all four analytical capabilities named above and still fail, if the institution around it hasn't provided the sponsorship and standing authority to act on what it produces.

It would be possible to go further here — to specify where the function sits on an org chart, who it reports to, how its outputs physically reach the rooms where capital and strategy decisions get made. That specification matters, but settling it isn't the goal here — the goal is to establish what any credible structure has to satisfy before it can be evaluated at all: that's the boundary this installment set for itself.

The discipline worth naming directly: an enterprise that accepts everything argued above still has a genuine design problem in front of it — what structure actually satisfies these requirements, staffed by whom, with what authority, producing what specific outputs. Different organizations, at different scales, with different risk profiles, could reasonably build different answers to that question and still be correctly responding to the same diagnosis. What they cannot do is skip the diagnosis and go straight to a structure, which is precisely the mistake the AI-procurement wave is currently encouraging across the industry.

The next installment introduces that framework — built to satisfy the requirements named here, and to be evaluated against them rather than taken on faith.

Works Cited

- [1] Carlile, P. R. (2002). *A Pragmatic View of Knowledge and Boundaries: Boundary Objects in New Product Development*. *Organization Science*, 13(4), 442–455.
- [2] Carlile, P. R. (2004). *Transferring, Translating, and Transforming: An Integrative Framework for Managing Knowledge Across Boundaries*. *Organization Science*, 15(5), 555–568.
- [3] Blondin, D., & Boin, A. (2020). *Cooperation in the Face of Transboundary Crisis: A Framework for Analysis*. *Perspectives on Public Management and Governance*, 3(3), 197–209.