

When a Single Agent Provably Cannot: An Elimination Criterion and a Cross-Domain Taxonomy of Coordination-Forcing Conditions

Hifzullah Celik
Ultradian
hifce@ultradianai.com

Abstract

When is more than one coordinating agent genuinely *required* for a task, as opposed to merely faster, cheaper, or more reliable with several? Most accounts of “why you need a team” describe a *degradation* gap: one worker would be slower, would run out of memory, would lack a skill. None of these establishes necessity: given enough time, memory, and competence, a single worker would suffice. This question has acquired sudden practical weight: as autonomous software agents scale, the things scale buys — context, patience, capability — are the very things that dissolve the degradation reasons, leaving open which coordination is structurally unavoidable.

This paper makes two principal contributions. First, a cross-domain taxonomy: applying a single elimination test across many domains by case-based induction over an illustrative survey of hundreds of real, named cases, we find that the coordination that survives is neither one thing nor unbounded but a small set — eight conditions, which we resolve more finely into ten — falling into five families, an organizing synthesis. A single locus of control is necessarily one body in one place at one time, one accountable identity, one center of interest and will, one knowledge-and-error state, and one continuous existence; each “one” is a wall when a task requires two. “Cannot” is proof relative to the idealization and, for the identity and interest families, relative to a humanly-imposed validity predicate; the modal force varies across families (Section 6.1), strongest for the physical conditions. Second, a sharp negative result: redundancy and fault-tolerance for *unreliable components* — the canonical “distributed systems need many nodes” reason — are *degradation*, not forcing, because they exist to compensate for components that crash or corrupt, which an idealized agent does not. We scope the taxonomy to safety-type forcing and flag availability-under-partition as a distinct, out-of-scope boundary.

The lens that makes both results possible is an elimination criterion: posit an idealized *maximal single agent* — unbounded patience, unbounded durability (never fails, never forgets), and full per-task capability, acting as one serial locus of control — and admit a condition as coordination-forcing only if this agent *still* cannot perform the task at all while some finite coordinated set can. Everything scale could fix is removed by construction; what remains is structural. The move of idealizing a single agent and reading off the residue has antecedents (Steiner’s task taxonomy; Brooks); our contribution is to take it to an *unbounded* limit and across non-social domains.

The taxonomy is induced from an illustrative, cross-domain case survey, and we subject it to a robustness check: a second application of the same criterion to a separately assembled corpus — sharing no cases or intermediate findings with the first — reproduced the same negative result and the same conditions up to resolution, and, once the two condition sets were reconciled, fell on the same family boundary (the two differed on how finely they cut within the families, and on one thin condition each surfaced that the other did not). Because both applications share the criterion and the investigator, this is closer to test–retest than to independent replication: it shows robustness to corpus, seeding, and analysis order, not mind-independence (Section 5). Finally, we position the criterion against the recent monotonicity

(CALM) account of avoidable coordination, showing the two questions are largely orthogonal: tasks that require coordination under monotonicity are often *not* forcing under elimination, and tasks that are forcing are classified coordination-free under monotonicity’s natural (additive) representation, the orthogonality being sharpest on the identity, interest, and existence families that its criterion does not condition on at all.

1 Introduction

Coordination is expensive. Meetings, approval chains, hand-offs, sign-offs, and cross-checks all consume effort that the underlying work does not, and organizations devote a substantial share of their resources to them. The same is now true, and measurably so, of multi-agent software systems, where the tokens spent on an orchestrator planning, routing, and reconciling sub-agent outputs can rival or exceed the tokens spent on the work itself. In both settings the operative assumption is that the coordination is *necessary*. Much of it may not be: recent empirical work finds that automatically generated multi-agent systems frequently underperform a single strong agent at up to an order of magnitude more cost [24], which is live evidence that much present-day multiplicity is a non-forcing optimization that underdelivers — and which sharpens the question this paper answers, of when coordination is instead structurally *forced*.

The difficulty is that “necessary” is rarely defined. When we say a task “needs a team,” we usually mean one of three things, all of which are claims about *degradation* rather than necessity:

- **It would be too slow for one.** A single worker would take too long. But slowness is relative to a deadline that, absent an external clock, can be met by waiting.
- **It would overwhelm one.** A single worker could not hold enough in memory, or process enough volume. But this is a limit of finite memory and throughput, not of singleness.
- **One could not do all of it well.** A single worker lacks some specialist skill. But this is a limit of one agent’s competence, not of the number of agents.

Each of these dissolves the moment we imagine a worker patient enough, capacious enough, and capable enough. They describe situations where multiplicity *helps*; none describes a situation where multiplicity is *required*. These limits are what the current generation of agentic systems is engineered to push back: longer context windows, persistent memory, higher reliability, broader capability. If the only reasons to coordinate were degradation reasons, then a sufficiently scaled single agent would eventually need no team at all.

So the sharp question is the residual one: *what tasks remain impossible for a single agent even after we remove slowness, forgetfulness, capacity limits, and skill gaps entirely?* What is left when scale is assumed to have already won?

This paper answers that question with a taxonomy and a negative result, derived through a single enabling lens.

The taxonomy (Section 3) is the primary finding: applying one elimination test across many domains by case-based induction, the conditions that survive are ten in number, organized into five families, where each family corresponds to a different respect in which a single serial locus is necessarily *one*: one body/place/time, one accountable identity, one interest/will, one knowledge-and-error state, one continuous existence. The conditions are the high-confidence result; the five-family organization is an explanatory synthesis.

The negative result (Section 4) is the second finding, and the least contestable: redundancy and fault-tolerance *for unreliable components* are not coordination-forcing. Triple-modular redundancy

and N-version programming mask the value and corruption faults of components; Paxos and Raft mask crash and omission faults; Byzantine fault-tolerant consensus masks arbitrary, including adversarial, faults — all to make an unreliable collection of real components behave like a reliable whole. An idealized, fully reliable agent suffers none of these, so the premise that motivates redundancy is removed and the apparent forcing dissolves. The genuine structure in that neighborhood does not vanish — it redistributes onto conditions the taxonomy already names (surviving the destruction of the locus; an error-decorrelated check; anti-collusion). We scope the claim carefully: it concerns reliability, and we set *availability under partition* aside as a distinct boundary the present taxonomy does not house (Section 4).

Both results are produced by the same enabling lens, the elimination criterion (Section 2): posit an idealized maximal single agent (unbounded patience, unbounded durability, full per-task capability, one serial locus of control) and admit a coordination-forcing condition only if this agent *still* fails while a finite coordinated set succeeds. Removing every degradation reason in the construction of the idealization guarantees that whatever survives is a structural fact about being a single locus, not a contingent fact about being a limited one. The criterion is an idealization *move* with clear antecedents — Steiner’s analysis of which tasks are bounded by structure rather than member capacity, and Brooks’s “the bearing of a child takes nine months, no matter how many women are assigned” — applied here at an unbounded limit and across non-social domains; we examine its standing, edge cases, and antecedents in full in Section 6.

The taxonomy is induced from an illustrative, cross-domain case survey, and we run a robustness check on it (Section 5). The same criterion was applied a second time over a separately assembled corpus of real, named cases from different domains, sharing no cases, intermediate work, or conclusions with the first. The two converged on the same negative result and on the conditions up to resolution, and — once their condition sets were reconciled — on the same family boundary, differing on within-family granularity and on one thin condition each surfaced that the other did not. Because both applications share the same criterion and investigator, this second pass is a sensitivity check — closer to test–retest than to independent replication — and it shows the structure is *robust to corpus, seeding, and analysis-order variation* rather than an artifact of one assembly or one path through the data. Section 5.3 sets out the stronger tests (a blind inter-rater statistic, a second operationalization, expert validation) that a future, auditable study would run.

Finally (Section 6), having presented the two contributions, we examine the criterion itself — its idealization’s edge cases, the differing modal force of “forced” across the families, and its antecedents — and then (Section 6 Related work) position it relative to organizational coordination theory, distributed-systems impossibility, and, in detail, the most recent and most closely related result: a monotonicity-based account of when coordination is avoidable. We show that the elimination criterion and the monotonicity criterion ask nearly orthogonal questions, and that the most natural reading of the monotonicity account’s own stated scope is that our taxonomy occupies precisely the territory it sets aside.

Contributions.

1. A **cross-domain taxonomy** of ten coordination-forcing conditions in five families, induced from an illustrative survey of real, named cases spanning physical, computational, biological, legal, financial, scientific, and operational settings — conditions in which being a single serial locus is *itself* the wall.
2. A **negative result**: redundancy for unreliable components is degradation, not forcing — stated explicitly, located within the taxonomy, and scoped to reliability (with availability-under-partition flagged as out of scope).

3. The **elimination criterion** as the enabling lens: defining coordination-requirement against an idealized maximal single agent, which separates structural necessity from degradation cleanly and in a way suited to the agentic era — an idealization move with antecedents (Steiner; Brooks) taken to a novel unbounded limit and cross-domain reach.
4. A **robustness check** on the taxonomy — a second application of the criterion over a separately assembled corpus — showing the structure is robust to corpus, seeding, and analyst variation. Because the two passes share a criterion and an investigator, this is a sensitivity signal (test–retest), not independent replication; Section 5.3 lays out the auditable tests that would carry it further.
5. A **positioning** that distinguishes the criterion from the monotonicity (CALM) account and from the prior camps — including the one prior *formal* required-cooperation account, whose causes the criterion subtracts off as degradation — with orthogonality established on the identity, interest, and existence families the monotonicity criterion does not condition on, and a conceded overlap on the epistemic family.

2 The elimination criterion

Before the taxonomy, the reader needs the exact test in hand, because every condition’s admission turns on it. We state it compactly here; its idealization edge cases, modal grades, and antecedents are examined after the results, in Section 6.

2.1 The maximal single agent

Let the *maximal single agent* be an actor granted three idealizations and held to one structural constraint.

- **Unbounded patience.** Time is free. The agent never tires, is never rushed, and may take arbitrarily long over any step. Any obstacle that is merely “too slow for one” is removed.
- **Unbounded durability.** The agent never fails and never forgets: perfect, unbounded memory and perfect reliability. Any obstacle of the form “one could not hold enough” or “one might make an error or crash” is removed.
- **Full per-task capability.** The agent performs any individual subtask as well as the best available specialist. Any obstacle of the form “one lacks a skill” is removed.
- **One serial locus of control.** The agent is one actor, one identity, taking one action at a time, in sequence. This is the *only* thing it is not granted relief from. It is one body, one will, one knowledge-state, one continuous existence.

The first three grants dissolve every degradation reason; the fourth is the property under test. One principle governs their reach: *the grants remove the agent’s endogenous limits (its own patience, memory, skill, and reliability) and nothing else.* They do not alter the *exogenous* facts of the world the agent acts in: the external clock, the medium that separates two locations, and the agent’s own finitude (that it can be destroyed by something outside itself). We idealize the agent to perfection and leave the world as it is; this is what keeps the idealization a model of an *agent* rather than the absence of a world, and two of the conditions below (B and J) are forcing precisely because they turn on an exogenous fact the grants do not reach. We defend this endogenous/exogenous line, and work through its edge cases, in Section 6.1.

2.2 The admission rule and the forcing band

A task is **coordination-forcing** if and only if the maximal single agent cannot complete it *at all*, while some finite set of coordinated agents can. Being merely slower, costlier, more error-prone, or memory-bound does not qualify, because the idealization has removed all of those as possibilities.

This places forcing tasks in a band bounded on both sides, *between what a single agent can do given enough time, memory, and skill (the degradation regime below) and what no finite set of agents can do (genuine impossibility above)*. Below the band, multiplicity is a convenience; above it, multiplicity does not help either; coordination-forcing tasks are exactly those in between — impossible for one, possible for several.

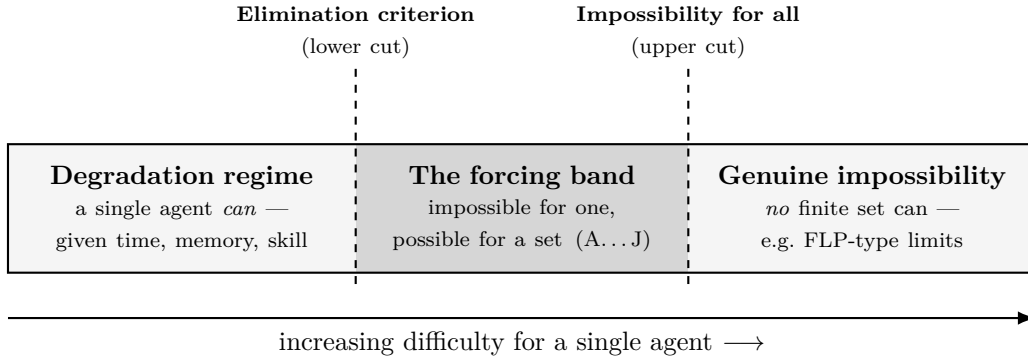


Figure 1: **The forcing band.** A task is *coordination-forcing* when it falls in the band between what a single agent can do given unbounded time, memory, and capability (the degradation regime, left) and what no finite set of agents can do (genuine impossibility, right). The elimination criterion is the lower cut: a condition is admitted only if an idealized *maximal single agent* — unbounded patience, unbounded durability, full per-task capability, one serial locus — still cannot perform the task while some finite coordinated set can. The ten conditions of Section 3 populate the band.

The band makes the criterion falsifiable in a specific, useful way. To *reject* a candidate condition, one need only exhibit a way for the maximal single agent to accomplish the task by sequencing — to “serialize” it — using its unbounded patience, memory, or capability. A condition is admitted only when no such serialization exists. This is the test we apply adversarially throughout: for every candidate, we actively try to defeat forcing by constructing a single-agent serialization, and we keep only the conditions for which the attempt provably fails. Two conditions (B and J) are admitted only under a particular reading of where the idealization’s edge lies; we flag them at the point of use and examine the reading in Section 6.1.

3 The taxonomy: ten conditions in five families

The conditions that survive the elimination test are not a flat list. They sort cleanly by *which* respect of singleness is the binding obstacle. A single serial locus is necessarily one in each of five respects, and each “one” generates its conditions, as the fan-out below sets out.

For each condition we give the *predicate* (the verifiable structural requirement), the *elimination argument* (why the maximal single agent provably fails), the reason it is *not mere degradation*, and *real, named exemplars* drawn from the survey. We lead with the conditions themselves, which are the primary finding; the five-family organization is the synthesis that orders them (Section 3.6).

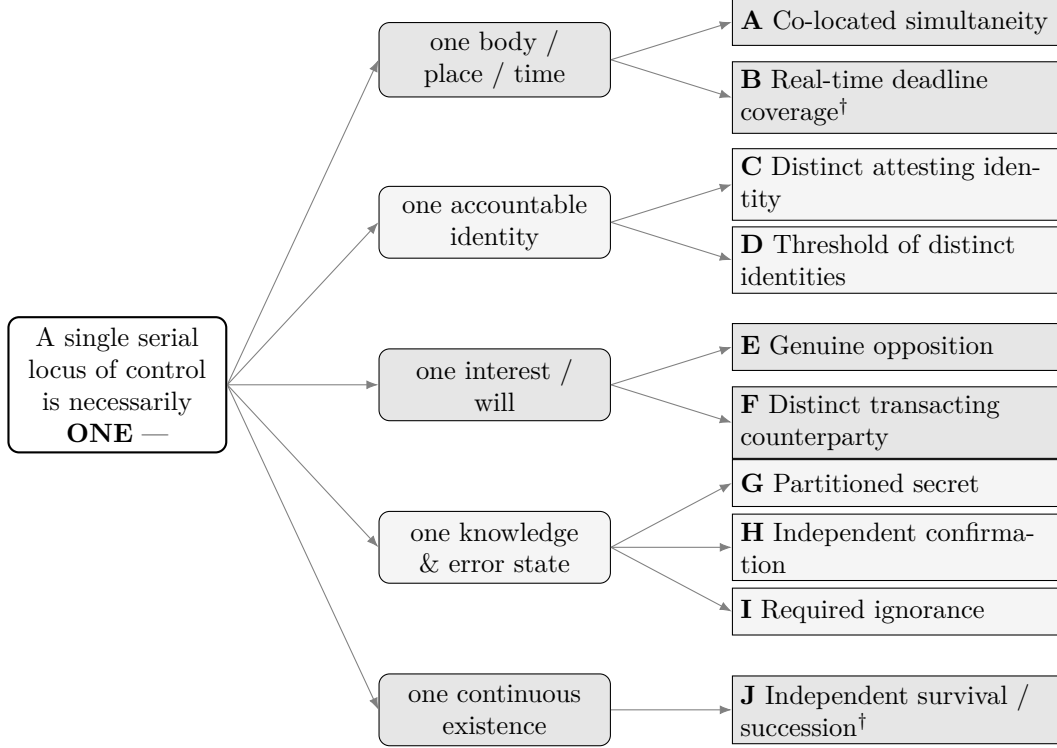


Figure 2: **The five singularities of a single locus of control, and the ten coordination-forcing conditions they generate.** Each family corresponds to a respect in which a single serial locus is necessarily *one* — one body/place/time, one accountable identity, one interest/will, one knowledge-and-error state, one continuous existence — and each “one” is a wall when a task requires two. The ten conditions are the high-confidence finding; the five-family axis is an organizing synthesis (Section 3.6, Section 7). [†] Conditions B and J are forcing under the idealization reading of Section 6.1.

3.1 Family I — Spatiotemporal singularity (*one body, one place, one time*)

A. Co-located simultaneity. *Predicate.* The task requires two or more physical actions to occur *at the same instant* at points separated by more than one body can span. *Elimination.* A serial locus performs one action at one place at a time. Two simultaneous, separated actions cannot be stitched together by sequencing, because the binding constraint is *coincidence in time*: unbounded patience cannot serialize an instant, perfect memory cannot recall it into being, and full skill cannot substitute for a second simultaneous effector. Any rigid jig or linkage that bridges the gap is itself a second locus, barred by definition. *Not degradation.* Time, memory, and skill are all irrelevant. A two-hand machine press fires its die only while both palm buttons — spaced beyond one hand-span — are held within the same brief window; the agent that releases one button to reach the other never actuates the press, no matter how patient. *Exemplars.* The two-key launch interlock of a Minuteman missile (keys placed beyond one operator’s reach, turned within seconds of each other); OSHA two-hand press controls; the baton exchange in a 4×100 m relay, completed within a moving takeover zone; a two-person patient lift; the spatially separated brass choirs of Berlioz’s *Requiem*.

B. Real-time deadline coverage (forcing under the Section 2 idealization reading; see Section 6.1). *Predicate.* Two or more time-critical demands arise at *separated* places, with deadlines set by an *external* clock that will not pause; one serial locus, attending one at a time, inevitably misses others. *Elimination.* Distinct from A: no single instant demands two actions, yet the agent still fails. While it services demand 1, demand 2’s externally imposed window expires. The

agent’s patience is unbounded, but the world’s clock is not the agent’s to extend (Section 6.1). *Not degradation.* This is the subtle boundary. Throughput *with no deadline* is degradation — serialize arbitrarily slowly and finish eventually. It becomes forcing *only* when a hard external deadline makes “attend A, then B” miss B absolutely — not relative to a faster alternative, but against a fixed window. *Exemplars.* Air-traffic control of converging aircraft in adjacent sectors; search-and-rescue against an avalanche survival window; the cockpit division into pilot-flying and pilot-monitoring; the simultaneous compressions, airway, and defibrillation of a resuscitation.

3.2 Family II — Identity / authority singularity (*one accountable identity*)

C. Distinct attesting identity (separation of duties). *Predicate.* Validity requires that the party who checks, witnesses, approves, or authorizes be a *different accountable identity* from the party who acted. *Elimination.* A single locus has exactly one accountable identity. However flawless its self-review — even after an unbounded wait, even with a deliberately adversarial fresh framing — what it signs is *self-attestation*, which the rule structurally refuses to count. It can change its mind, its method, and its mood; it cannot change *whose single signature* the approval bears. *Not degradation.* The wall is the validity predicate “approver $\neq$ actor,” not the quality of the review. Perfect review quality does not satisfy a rule that is about identity, not about quality. *Exemplars.* The surgical pre-incision time-out; financial maker-checker (“four-eyes”) control [14], whose failure is the archetype of a single trader who was also his own back-office check; the maxim that no one may be judge in their own cause; code review by a non-author; the ceremonial role separation in a DNSSEC key-signing event.

D. Threshold of distinct identities (quorum). *Predicate.* Validity is defined as a *count* of $k \geq 2$ distinct identities, each casting their own act, vote, or consent. *Elimination.* One identity acting repeatedly over unbounded time produces many *acts* but still tallies as *one identity*. The deficit is in identity-cardinality, which neither time nor memory can manufacture. A jury “voted” by one agent occupying every seat is a tally of one. *Not degradation, and distinct from C.* A jury is pure threshold: every member occupies the same symmetric role, validity is a count, and there is no actor-versus-checker asymmetry. A single notary attesting one signer is pure separation-of-duties: relational asymmetry with no threshold. Each occurs without the other. What makes a *count of distinct identities* the validity predicate is itself an institutional rule — a constitutive “counts-as” convention under which the assent of a plural body is the act of distinct parties, not one party repeated (Searle [15]; Gilbert [16]) — so the wall in this family is forcing relative to a humanly-imposed predicate, in contrast to the physical wall of co-located simultaneity (A). We return to this difference of modal force in Section 6.1. *Exemplars.* A jury (twelve at common law, fewer in many jurisdictions); the attesting witnesses a formal will requires; bicameral passage plus presentment to the executive; constitutional amendment by separated supermajorities.

3.3 Family III — Interest / will singularity (*one center of aims and holdings*)

E. Genuine opposition (adversarial contest). *Predicate.* The outcome is produced only by two or more parties striving for *genuinely opposed*, independently grounded aims, with moves the other does not control. *Elimination.* A single will cannot hold two independently grounded aims and act to maximize each against the other; it can only *choose* an outcome, not have it *contested out of* a real opposition. (A single system playing both sides — self-play, or one author composing a game’s lines for both colors — does not refute this: there is one top-level aim, served whichever side “wins,” so the outcome is chosen, not contested out of two stakes neither side controls. Genuine opposition requires two centers of interest, each wanting to win and not indifferent to which side

does.) The result of a genuine contest is information that exists only because neither side controls the other. *Not degradation, and distinct from F.* A chess game or a tug-of-war is pure opposition: opposed striving with no transfer of property and no price. A willing-buyer/willing-seller sale is pure transaction: a cooperative transfer with no opposition. They double-dissociate. *Exemplars.* A tug-of-war; a penalty shootout; the bowler–batsman contest in cricket; adversarial cross-examination (secured, in criminal trials, by the right of confrontation, and more broadly by the adversarial structure of trial); adversarial collaboration in science.

F. Distinct transacting counterparty. *Predicate.* A transfer of value, a discovered price, or an enforceable bilateral obligation requires two genuinely distinct property-owners whose valuations differ (distinct reservation prices) so that a price can exist between them. *Elimination.* With one owner on both sides, ownership before equals ownership after: no title moves, no consideration flows, and a price can only be *discovered between* two distinct reservation values — with one owner there is no second value, so the bargaining range collapses to a point and there is nothing to discover. A self-sale is a structural no-op, eliminated rather than slowed. (On consolidation, a transfer between two accounts of a single owner nets to zero: it may be booked, but it creates no enforceable obligation against a distinct party.) *Not degradation.* A debt owed to oneself, or a price “discovered” against oneself, is not a weak transaction; it is a non-transaction. *Exemplars.* A bilateral sale or contract; the assignment or novation of a contract; a benchmark price-discovery auction such as the daily gold fixing, in which submitted orders from distinct participants discover a single reference price. (Auctions and collective bargaining combine F with the genuine opposition of E; the cleanly dissociating pair is a willing-buyer/willing-seller sale for F against a contest such as chess for E.)

3.4 Family IV — Knowledge singularity (*one knowledge-and-error state*)

G. Partitioned secret (no single holder). *Predicate.* The guarantee is that *no single locus ever holds enough to reconstruct or know the whole*, while the whole must remain functional — signable, decryptable, computable — through the combination of separately held parts. *Elimination.* A single locus that can produce the result necessarily holds every part at once, which is exactly the forbidden state. More memory makes this *worse*, not better. *Not degradation, and distinct from I.* The wall is a prohibition on a single *knower of the whole*, orthogonal to volume, time, skill, and reliability. It double-dissociates from required ignorance (I) across real multi-holder cases: dual-control key custody — where a key is split so that no single custodian can ever assemble it, yet no custodian is required to be *ignorant* of anything beyond the other’s share — is partitioned holding *without* required ignorance; allocation concealment (I) is required ignorance *without* a partitioned secret: one party must act while genuinely not knowing the assignment, and no whole is split across functional holders. (Secure multiparty computation is *not* a clean witness for the first direction: its defining guarantee is input privacy, so it instantiates *both* G and I at once.) Against the single idealized agent specifically the distinction narrows, since a lone locus able to produce the result would have to hold the whole, which is also the forbidden non-holding state of I; this is *why* G is the most contested condition in the set (Section 3.6). We retain it on the strength of the multi-holder dissociation and its distinct positive structure: it requires two or more functional-but-non-whole-holding parties, where I requires only one ignorant party. *Exemplars.* Threshold cryptography, in which a signature or decryption is produced from separately held shares while the key is *never reassembled anywhere* — the genuinely forcing form (plain secret-sharing [13] that is later *reconstructed* assembles the whole in one place at the moment of reconstruction, and is the policy/serializable boundary case the criterion correctly screens off); split-knowledge key loading in a hardware security module; secure multiparty computation, in which several parties jointly

compute a result while no party learns anything about another’s input beyond what the result itself reveals.

H. Independent confirmation (error decorrelation). *Predicate.* The task requires a confirmation whose *errors are uncorrelated* with the actor’s. *Elimination.* A single, never-forgetting mind is one error source; every re-examination inherits its one set of biases and choices, and it cannot approach its own conclusion fresh — perfect memory makes this worse, not better. It can never be the independent second source the check demands. *Not degradation, and distinct from C and I.* A double-check of a high-alert medication can be performed by a second nurse who is *not* an accountable approver (confirmation without separation-of-duties). Independent replication is satisfied even when both laboratories know everything (confirmation without required ignorance). Conversely, a four-eyes approval can be satisfied by a distinct identity whose training is highly *correlated* (separation-of-duties without confirmation). *Exemplars.* Independent replication of a scientific result; the cross-checking of an extraordinary physics measurement by two separately built detectors; the independent second check of a high-alert medication; two-source corroboration in journalism.

I. Required ignorance (blinding). *Predicate.* The task requires a party in a genuine state of *not-knowing* something. *Elimination.* A never-forgetting mind cannot un-know what it knows. If the single agent set up the arrangement — the allocation, the assignment — it knows it, and the required not-knowing state is unreachable. Blinding needs a party who genuinely lacks the information while still acting. *Not degradation, and distinct from C and H.* Allocation concealment in a blinded trial is satisfied by *any one party* who does not know the assignment — no second accountable identity, and no second checker. *Exemplars.* Allocation concealment in a double-blind randomized trial; placebo administration under single-blind; outcome adjudication by a committee kept blind to assignment.

3.5 Family V — Existential singularity (*one continuous existence*)

J. Independent survival / succession (forcing under the Section 2 idealization reading; see Section 6.1). *Predicate.* The capability must persist *through the destruction or incapacitation of the acting locus itself* — the success condition is located after, or despite, that locus’s removal. *Elimination.* The idealization grants “never fails / never forgets” but not invulnerability to external destruction (Section 6.1). A task whose success is defined as “X happens after I cease to exist or am incapacitated” — a will taking effect after the testator’s death; control continuing after a pilot is incapacitated — cannot be done by that single locus, because doing X requires existing past the stipulated end. A pre-armed automaton or a deputy *is* a distinct surviving locus, which concedes plurality. *Not degradation.* This is not a reliability problem. The agent is perfectly reliable while it exists; the task simply extends past its existence. *Exemplars.* The minimum-crew rule that defends against single-pilot incapacitation; a will or post-mortem trust; a line of succession; standing fail-deadly arrangements designed to act after their principals are gone.

3.6 The organizing structure

The ten conditions partition by which singularity of the lone locus is the obstacle, and the partition yields a one-line generalization:

Multiplicity is forced exactly when a task requires two values of a property of which a single locus of control can hold only one — one occupancy, one identity, one interest, one knowledge-state, one existence.

This five-family organization, and the generalization that summarizes it, are an *explanatory synthesis*: in the inductive work the conditions emerged as a flat set, and the families were read off the *reasons* that recurred in grouping them (“a single physical body can occupy only one location”; “a single identity’s self-review cannot count as independent”; “a single will cannot strive to defeat itself”; “a single never-forgetting mind cannot un-know”; “the very entity that is removed cannot act past its removal”). The conditions are the result; the family axis is the model that orders them. The families also differ in the *modal force* of their “can hold only one” (physical, epistemic, and institutional), a distinction that bears on the mind-independence question and which we take up in Section 6.1.

A coarser reading is defensible. The closest pairs (E/F, H/I) each share a family root: collapsing the two cleanly-paired soft spots ($E/F \rightarrow \text{one}$, $H/I \rightarrow \text{one}$) yields eight conditions in five families; a reader who additionally folds the contested partitioned-secret condition (G) into required ignorance (I) reaches seven. We report the finer ten because the paired conditions cleanly double-dissociate (Section 5), but we flag eight-in-five — and, more conservatively, seven — as legitimate alternatives.

4 The negative result: redundancy for unreliable components is degradation, not forcing

The most consequential thing the elimination criterion *rejects* is the family of reasons a distributed-systems practitioner reaches for first: redundancy and fault-tolerance. These mechanisms mask faults internal to an ensemble of real components, making an unreliable collection behave like a reliable whole: triple-modular redundancy and Avizienis’s N-version programming [22] mask the value and corruption faults of components; Lamport’s Paxos [20] and Ongaro and Ousterhout’s Raft [21] mask crash and omission faults; Byzantine fault-tolerant agreement [10] masks arbitrary, including adversarial, faults.

Under the elimination criterion this motivation dissolves, because the maximal single agent *never fails and is fully reliable*. A single perfect locus has no faulty peers to outvote and nothing to reach consensus *about* among parts of itself. The premise that motivates redundancy has been removed by construction. Reliability-through-redundancy is therefore a *degradation* concern: it closes the gap between a real, unreliable system and the reliable ideal — the gap the idealization assumes already closed.

This is not a claim that fault-tolerance is unimportant; it is a claim about *which axis it lives on*. Reliability is real and often dominant in practice. But it is not what makes a task require multiplicity *in principle*; it is what makes a *realization* of the task survive imperfect parts. Conflating the two is the error the criterion is designed to prevent.

The genuine structure that hides inside “we used multiple nodes” does not vanish; it redistributes onto conditions the taxonomy already names:

- *Keep acting through the destruction of the acting locus* is condition J (independent survival) — forcing because the success criterion refers to a post-destruction state, not because components are unreliable.
- A check whose *errors are uncorrelated* with the producer’s is condition H (independent confirmation) — forcing because a single never-forgetting mind is one error source, not because of crash-tolerance.
- *No single party may unilaterally act* — the real structure of two-person rules and anti-collusion controls — is conditions C/D (distinct attesting identity; threshold), forcing on identity grounds.

This bears directly on fleets of near-identical reasoning agents: replicating one model and polling the copies masks no fault the idealization left standing — identical reliable units voting tally as one — and any genuine check among them is forcing only insofar as it supplies condition H’s error-decorrelation, which copies of a single generator, sharing an error manifold, cannot supply from within. The decorrelated confirmer H demands must then come from outside that generator — a different model class, instrument measurement, or a human — not from replication.

Two residues are *not* removed by per-component reliability, and we name both rather than let them hide under “redundancy.”

Adversarial compromise. Reliability is not incorruptibility. A perfectly reliable component can still be *suborned* by an external adversary, and a requirement to tolerate up to some number of compromised participants — the modern reading of Byzantine fault-tolerance — does not dissolve under “never fails.” But this is not a new condition: a single locus, if compromised, is compromised in whole, so the requirement reduces to “no single party may unilaterally produce the result” (conditions C/D). The tolerated-compromise bound sharpens the threshold into an adversary-dependent super-majority ($n \geq 3f + 1$ to tolerate f Byzantine participants) rather than a bare count, backed where needed by error-decorrelation (H) and partitioned holding (G). Adversarial fault-tolerance is forcing, but on identity and knowledge grounds the taxonomy already supplies — not on reliability grounds.

Availability under partition — a deliberate scope boundary. The second residue is the one the idealization does *not* reach (Section 2.1, Section 6.1): the medium between two locations is part of the world, and per-agent reliability does not make it reliable. A task defined as “remain available (continuing to serve operations) at two locations whose connecting link may be severed” is not satisfiable by one locus, which sits at one place; and, by the CAP theorem (Gilbert and Lynch) [19], it is not jointly satisfiable with strong consistency even by *two* perfectly reliable loci under partition. This is fault-tolerance-adjacent forcing, and it is distinct from everything above: condition J concerns a locus that *dies*, whereas here the locus is alive but *cut off*; and the conditions of Section 3 are *safety-type* (no incorrect outcome is produced), cleanly so for the identity, interest, knowledge, and existence families, with real-time deadline coverage (B) sitting at the timeliness edge of that class, whereas availability is a *liveness* property (an outcome is still produced). The present taxonomy characterizes *safety-type forcing* — the conditions under which a single locus cannot produce a *correct* result at all — and it does not house liveness-under-partition. We mark this as the taxonomy’s most significant boundary, not as a covered case: neither inductive study surfaced it, because both corpora were assembled around safety-type forcing, and we decline to invent a condition the evidence did not produce (Section 5). A full account unifying the safety conditions here with a liveness/availability axis is the most important direction the work leaves open (Section 8).

The lesson generalizes within that scope: most “distributed systems need many nodes” arguments are about *unreliable components*, which the idealization removes; when a genuinely structural *safety* requirement is present, it survives the idealization and appears as one of the named conditions. Both studies bear this out (Section 5): neither yielded a separate “redundancy” condition, and under the criterion the structure of fault-tolerance cases belongs to J, H, and C/D.

5 The case survey, and a robustness check

The evidence for the taxonomy is an *illustrative, cross-domain case survey*: a large collection of real, named cases, each subjected to the elimination test, from which the conditions are induced. The numbers reported below — how many cases were assembled, how many survived as forcing, how

they distribute across conditions and domains — describe the *range and spread of that survey*, the cases the conditions were induced from and mapped against. They are not an audited benchmark: there is no released corpus, no published codebook, and no per-case expert validation, so the argument does not rest on the counts as proof. It rests on the conditions themselves — each a structural argument anyone can check against the named exemplars — with the survey showing the breadth over which they recur.

The taxonomy is an inductive claim, and the standing risk with one is that a single investigator imposes a preferred structure on the cases. To probe that, the same criterion was applied a second time to a separately assembled corpus, sharing no cases or intermediate findings with the first, and the resulting condition sets were reconciled against those of the first. Because both passes share the criterion and the investigator, this is a *robustness / sensitivity check* — a second application of one criterion, closer to test–retest than to independent replication — not the agreement of two independent teams. It is a useful signal, not a headline; the negative result and the conceptual contribution carry the paper, and the second corpus tells us the structure does not depend on which cases were collected or in what order they were read.

5.1 The method, stated plainly

The conditions were derived by cross-domain, case-based induction: assembling large collections of real, named cases in which coordination appears to be required; subjecting each case to the elimination test adversarially (actively attempting to show that a single patient, capacious, capable agent could serialize it); discarding the cases that serialize as degradation; and grouping the structural requirements of the surviving cases into conditions. The criterion’s falsifiability (Section 2.2) is what makes this disciplined rather than impressionistic: a case is retained as forcing only when a single-agent serialization is shown not to exist, and a condition is retained only when removing it leaves some case unexplained.

Three properties were checked for every retained condition:

- **Coverage.** Every forcing case is accounted for by the ten conditions — the fault-tolerance cases, which the negative result reclassifies as degradation, through their redistribution to the survival, confirmation, and identity conditions (J, H, and C/D; Section 4) rather than through a separate redundancy condition.
- **Minimality and independence.** Each condition is the *sole* forcing reason in a healthy fraction of its cases, so no condition is a mere combination of others. Every retained pair was tested for *double-dissociation* — exhibiting a real case with X-but-not-Y and another with Y-but-not-X — and pairs that failed to dissociate were collapsed rather than kept.
- **Generativity.** The conditions were tested against 84 fresh cases assembled from domains not used to derive them: polytempo music, molecular machinery, diplomacy, tabletop games, space-flight, and others. Of these, 70 were independently judged forcing; of those 70, 63 (90%) mapped cleanly onto the existing conditions, 4 mapped with a stretch, and 3 to “none/new.” On review the three apparent gaps yielded no validated new condition — two resolved to combinations of existing conditions, and one (a single spacewalk case turning on a body’s inability to reach its own exterior) remains a flagged, under-evidenced candidate we do not admit — it appeared in no corpus case and only once across the fresh set, failing the admission bar that the ten clear: a condition must recur in the corpus (not only in out-of-sample cases) and separate under independent grouping, which even the thinnest admitted conditions do. The exercise stressed all conditions and generated, but did not confirm, any eleventh.

No specialized vocabulary is needed to state any of this: it is case-based induction, adversarial falsification of each case, and out-of-sample testing.

The cases were assembled and grounded, and the conditions induced, with substantial large-language-model assistance. Each condition stands as an elimination argument the reader can check by inspection against its named exemplars; it is those arguments, not the provenance of the cases, that carry the claims.

5.2 The second corpus: a sensitivity check

The check is that *the same criterion, applied a second time over a different corpus and sharing no findings with the first, reached the same structure*. The primary survey assembled 263 real, named cases across 17 domains and, after adversarial re-examination, retained 193 as forcing; their structural requirements were grouped by several analyses working in different orders, which agreed on the same families. A second survey independently assembled a comparably sized corpus — 253 cases across 13 domains, of which 141 were retained as forcing — sharing no cases, intermediate work, or conclusions with it; reconciling the two condition sets afterward, the two fall on the same family boundary, which neither survey was told to look for.

That dividing line deserves stating in its own terms; neither survey was told to look for it, and it emerged only when the two condition sets were reconciled. The families split into two groups by a single test: *does the obstacle dissolve if the single agent is granted multiple bodies or effectors?* For the spatiotemporal family it does — one locus is one-thing-at-a-time, and replicated simultaneous capacity removes the obstruction (while conceding multiplicity). For the identity, interest, knowledge, and existence families it does not — more effectors share the one locus’s identity, stake, knowledge-state, and fate. One study framed this as a two-way split, the other as five families; it is the *same boundary*.

The differences between the two were differences of resolution and of one thin condition, not of the boundary. One resolved the identity, interest, and knowledge families more finely, surfacing the partitioned-secret (G) and required-ignorance (I) conditions as distinct where the other left them merged; the other resolved the spatiotemporal family more finely, and surfaced one narrow condition (a requirement for mutually exclusive physical states) that the first did not reproduce and we do not include in the ten. The two condition sets thus differ by a small symmetric difference in both directions. Neither reversed a condition the other admitted, the reconciled condition sets fell on the same family boundary, and both produced the negative result of Section 4 — the fault-tolerance drop — which, unlike the family boundary, is not entailed by the criterion’s wording. That the negative result recurred is the most informative part of the check, for the reason the next paragraph makes precise.

Because both passes share the criterion, they are not all equally surprising, and reading the check correctly means weighting its strands. The *family boundary* (does the obstacle dissolve with more effectors?) is close to a consequence of testing a single-locus constraint, so its reappearance is expected. The *negative result* is contingent — nothing in the criterion’s wording anticipates it — yet it appeared in both corpora; that is the strongest strand. The *condition set* is stable within each pass but differs between them in resolution, as above. What the check cannot do is establish mind-independence: a shared criterion tends to produce a shared result, so a second pass under it speaks to robustness, not to whether the criterion itself imposes the shape (Section 5.3). What keeps the result more than a restatement of the criterion is that the criterion is underspecified about its output — “admit what survives elimination,” not “expect a small closed set” — yet a small, largely shared set plus an unanticipated negative result emerged twice.

Cond.	Condition (short)	Family	Corpus cases	Domains	Sole-forcing reason	Out-of-sample mapped
A	Co-located simultaneity	Spatiotemporal	24	7	yes	25/70
B	Real-time deadline coverage [†]	Spatiotemporal	16	3	yes	5/70
C	Distinct attesting identity	Identity	42	10	yes	17/70
D	Threshold of distinct identities	Identity	16	4	yes	16/70
E	Genuine opposition	Interest/will	17	5	yes	8/70
F	Distinct transacting counterparty	Interest/will	13	2	yes	10/70
G	Partitioned secret	Knowledge	17	7	yes	3/70
H	Independent confirmation	Knowledge	17	5	yes	9/70
I	Required ignorance	Knowledge	7	2	yes	11/70
J	Independent survival / succession [†]	Existence	17	7	yes	1/70

Table 1: **Survey scope and out-of-sample mapping.** The counts report the *range and spread of the illustrative case survey* the conditions were induced from and mapped against — not an audited dataset; there is no released corpus or per-case expert validation (Section 5). The per-condition “Corpus cases” and “Out-of-sample mapped” columns report each condition’s illustrative spread as per-condition hits, not an exhaustive partition: a case can map to more than one condition, so the columns need not sum to the 193 retained or the 70 fresh-forcing cases. Each condition is illustrated by real, named cases verified against the elimination test; “domains” counts the distinct settings in which it recurs. “Sole-forcing reason = yes” means the condition is the only forcing reason in a healthy fraction of its cases, the basis of the double-dissociation argument (Section 5.1). The second-corpus agreement (Section 5.2) is a sensitivity check: a separately assembled corpus, run through the same criterion, reproduced the same negative result and the same conditions up to resolution, and — after the two condition sets were reconciled — fell on the same family boundary. Out-of-sample columns report how 84 fresh cases in new domains mapped onto the conditions derived without them. *Aggregate (survey scope).* Primary survey: 263 cases across 17 domains, 193 retained as forcing (all retained cases accounted for by the ten conditions, fault-tolerance cases via the Section 4 redistribution); second survey (sensitivity check): 253 cases across 13 domains, 141 retained, assembled without shared cases or findings — reproduced the same negative result, agreed on the conditions up to resolution, and, once the two condition sets were reconciled, fell on the same family boundary. [†] Forcing under the idealization reading, Section 6.1.

5.3 What it does not establish, and what would strengthen it

The sensitivity check speaks to the *structure*, not to every case: it does not certify each individual case attribution — some may be wrong without disturbing the conditions they illustrate — and it does not turn the survey into an audited benchmark or the five-family axis into an independent finding. The confidence is in the over-determined shape, not in any single case or in the organizing diagram. Three things would convert this illustrative survey into reproducible, mind-independence-grade evidence: *blind inter-rater agreement* (a reported agreement statistic on a held-out set classified independently by separate parties, rather than a prose report of agreement); a *second operationalization*, a team that defines “forcing” differently and still converges, which would test the structure against the criterion rather than under it; and *domain-expert validation* of a sample of case attributions. These are the natural next steps for a future, auditable study.

6 The criterion, examined; and related work

The two contributions above stand on the elimination test, but the test itself invites scrutiny: its idealization has edges, its “forced” has more than one modal grade, and its central move has antecedents. We examine those here, having first put the results in hand, and then position the

criterion against the established and recent literature.

6.1 The criterion examined: edge cases, modal grades, and antecedents

The endogenous/exogenous line, and the two edge cases it creates. The grants of Section 2.1 remove the agent’s endogenous limits and leave the world untouched. An idealization that also rewrote the world’s clock, abolished the distance between places, or conferred immortality would no longer be modeling a maximally capable agent — it would be modeling the absence of a world, and “can a thing with no world do it alone” is not a question about coordination. Holding that line fixes three edges of the idealization that the taxonomy relies on; they are consequences of the principle, not separate stipulations:

- *The world’s clock (the boundary for condition B).* Unbounded patience lets the agent wait as long as it likes; it does not let the agent make *external* deadlines wait. When two time-critical demands arise at separated places under a hard external clock, the agent’s own patience is irrelevant — while it services one, the other’s externally fixed window expires. This is what makes real-time deadline coverage (B) forcing, and it is the precise sense in which B is “forcing under the idealization reading”: a stricter idealization that let the agent pause the world’s clock would demote it.
- *The agent’s finitude (the boundary for condition J).* The agent is perfectly reliable *while it exists*; “never fails” is reliability, not immortality. It is not immune to being destroyed or incapacitated by something outside itself, so a task whose success is defined as something happening *after* the acting locus ceases to exist is not reachable by that locus. This is what makes independent survival (J) forcing, and the sense in which J too depends on the reading: an idealization that conferred invulnerability to external destruction would demote it.
- *The medium between locations (the Section 4 scope boundary).* The agent is one body in one place; the channel that connects any two places is part of the world, not part of the agent, and the grants do not make that channel reliable or instantaneous. This is the edge the negative result of Section 4 must reckon with: per-agent reliability does not buy a reliable network, and tasks defined by *availability across a partitionable medium* therefore sit at a boundary we treat explicitly (Section 4) rather than fold in.

These are the principal interpretive commitments of the work. A stricter idealization that rewrote any of them would demote the corresponding conditions (B, J); we have argued that the endogenous/exogenous line is the principled place to stand — it is what keeps the idealization a model of an agent — but it is a choice, and Section 7 records it as the work’s main interpretive dependency.

The modal grade of “forced.” “Can hold only one” does not have a single modal force across the families, and the difference matters for how strongly “forced” should be read. One body/place/time (A, B) and one continuous existence (J) are physical or temporal necessities; one knowledge-and-error state (G, H, I) is epistemic; one accountable identity and one interest/will (C, D, E, F) are *institutional*, forcing relative to a humanly-imposed validity predicate (a constitutive rule, Searle; a plural subject’s joint commitment, Gilbert), not a mind-independent constraint. The taxonomy holds across all three grades, but “forced” is stronger in the physical family than in the institutional one.

The criterion’s antecedents. The elimination *move* itself — subtract off what a capable individual could do and ask what remains — is not new. Steiner’s task typology already distinguishes

additive, disjunctive, conjunctive, and divisible tasks and isolates the cases where group output is bounded by task structure rather than by member capacity; his decomposition of group performance into potential productivity minus process loss is, in effect, the degradation axis we subtract off. Brooks’s observation that “the bearing of a child takes nine months, no matter how many women are assigned” is the folk statement of an elimination argument. We differ from this lineage in two ways that matter: we idealize to an *unbounded* single agent rather than to the most capable available member, which sharpens “cannot” from “cannot efficiently” to “cannot at all”; and we range across non-social domains (cryptography, law, aviation, clinical method) that small-group psychology does not reach, which is what yields a closed cross-family set rather than a typology of group tasks. The move is an inherited lens; its unbounded idealization and its cross-domain reach are what is new.

6.2 Related work and positioning

Three established traditions ask adjacent questions, and a fourth, very recent result asks the most closely related one. We take each in turn and locate the elimination criterion against all four.

6.2.1 Three prior camps, and a near antecedent

Organizational and coordination theory. Thompson’s classic typology [1] classifies tasks by their pattern of interdependence (pooled, sequential, reciprocal) and observes that reciprocal interdependence is the costliest to coordinate. Malone and Crowston’s interdisciplinary program [2] defines coordination as the management of dependencies among activities and sets out to characterize dependency types and the mechanisms that manage them. Puranam and colleagues [3] sharpen the picture by distinguishing task interdependence from *epistemic* interdependence — what each agent must know about what others know. This tradition is indispensable, but it *presupposes* plurality: it asks how to organize many agents and when interdependence is costly, not whether a single idealized agent could have done the task alone. Its question begins where ours ends. One connection we draw rather than disown: Puranam’s epistemic interdependence is the closest prior analog to our knowledge family (G, H, I), and we read those conditions as a structural, elimination-grounded reconstruction of part of what that line identifies — not as its replacement.

Group productivity and the additivity of tasks. The group-productivity tradition (Steiner; Brooks) [17, 18] is the near antecedent of the criterion’s elimination move, examined in Section 6.1: it supplies the degradation axis we subtract off and the folk form of the elimination argument, while differing from our account in idealizing to an available member rather than an unbounded agent and in staying within small-group psychology rather than ranging across cryptography, law, aviation, and clinical method.

Distributed-systems impossibility. A second tradition asks what is solvable *inside a faulty, asynchronous model*. The impossibility of guaranteed-terminating consensus in an asynchronous system with even one crash [9] — whose engine is the indistinguishability of a slow process from a failed one — the Byzantine generals problem [10], and the topological characterization of asynchronous computability [11] all establish hard limits relative to *failure and asynchrony*. These are impossibility results, not convenience gaps: they sit *above the forcing band relative to the faulty, asynchronous model they assume* — no finite set escapes them within that model. But that placement is model-relative, and every one of these results turns on the possibility of failure (FLP on a crash that cannot be distinguished from slowness; the wait-free obstructions on processes that can crash). The elimination criterion removes per-locus failure, so against the failure-free idealization these results *dissolve*, which is why they are not forcing conditions (Section 4). What the idealization

does *not* remove is a property of the world that failure-removal cannot touch, and it is not the *asynchrony* of the medium (FLP assumes an untimed but reliable channel) but its *partitionability*: the connecting link may be severed and lose messages. That residue is the availability-under-partition boundary we house as a scope edge rather than a condition (Section 4); it is driven by partition, a reliability property of the medium, not by timing. So the tradition is complementary in two ways: where its results turn on failure or pure asynchrony, the idealization dissolves them; the genuinely irreducible residue is partition of the medium, which we treat explicitly. It characterizes the cost of imperfection and asynchrony; we characterize the residue that per-agent perfection cannot remove.

Formal required cooperation in planning. Closest in spirit to our question is the one prior *formal* account of when a single agent is insufficient: Zhang, Sreedharan, and Kambhampati’s analysis of *required cooperation* (RC) in multi-agent planning [23], which asks, in a bounded SAS+ planning formalism, exactly when a goal is single-agent-unsolvable. It is the natural source of the objection “isn’t this just required cooperation in planning?”, and answering it sharpens our contribution rather than threatening it. Their RC causes are of precisely two kinds — agent *heterogeneity* (an agent lacking some other agent’s action, domain, or variable capabilities) and *state-space structure* (a non-traversable causal graph, or causal loops the agent cannot break) — and they show that when none of these holds the problem is single-agent-solvable. Tellingly, they observe that such problems are typically *transformer-agent solvable*: a single virtual agent endowed with the union of all agents’ capabilities, able to traverse the combined state space, succeeds. That is exactly the maximal-single-agent move of Section 2 run in reverse, and it shows that their RC causes are the *bounded-capability and state-structure limits the elimination criterion subtracts off as degradation, not forcing*: grant the single agent full per-task capability and unbounded patience and every transformer-agent-solvable RC problem dissolves. Accordingly that account enumerates *none* of the structural singularities the present taxonomy turns on (identity, will, knowledge, existence) because none of them is reducible to a missing capability or an untraversable state; they are walls a single locus hits even with every capability combined. The one prior formal required-cooperation result is thus not a competitor but a confirmation: it formalizes, within planning, exactly the regime our criterion places *below* the forcing band.

6.2.2 The near neighbor: the monotonicity (CALM) account

The nearest neighbor defines *avoidable* coordination through monotonicity [4]. Building on the CALM principle — that a computation admits a coordination-free implementation if and only if its specification is monotone [5, 6, 7] — this account maps Thompson’s interdependence types onto the monotonicity criterion and derives a decision rule: coordination is required for correctness exactly when extending the execution history can *invalidate previously admissible conclusions* (non-monotonicity). It validates the rule on enterprise workflow corpora and occupational task databases, and illustrates it in multi-agent simulations.

This is the same *genre* — a cross-domain account of when coordination is required — so the natural question is whether the present taxonomy is “just monotonicity restated.” It is not. We give the argument in order of strength, leading with what a reader can check directly and treating the near neighbor’s own scope statements as corroboration rather than as the load-bearing step.

Different criterion object. The monotonicity account does not idealize a single agent at all. Its execution model is multi-agent from the outset — a set of agents under a happened-before order — and “coordination-free” means that independent, asynchronously executing agents are *never blocked to prevent future inconsistency*. The question it answers is: *can independent executors run without synchronizing and still stay consistent?* Ours is the opposite idealization: *imagine one flawless agent with unbounded time, memory, and capability — what does it still fail at?* The difference is

not cosmetic. The monotonicity criterion is a property of an *outcome lattice* (how outputs combine as histories grow) and it does not condition on *who* acts. A lattice can of course carry identity-tagged elements (a grow-only set of votes [8] is the standard example), but the criterion has no way to *require* that the approver differ from the actor, that an outcome be contested out of two opposed wills, or that a capability persist past the actor’s destruction: it asks only whether extending the history can invalidate prior conclusions. The identity, interest, and existence families (C, D, E, F, J) are therefore not what it classifies; their forcing requirements are not the kind of thing its question can pose. This is the core of the distinction and it depends on no quotation.

Double-dissociation, where the criteria can both speak. Where they overlap, they cut differently, and the dissociation is checkable. In one direction: the non-monotone class in that account is, in practice, dominated by *shared-finite-resource allocation* (dividing a fixed budget, headcount, inventory, or credit pool), non-monotone because independent allocators could over-commit the pool. But a shared-budget allocation is *not forcing* under elimination: a single maximal agent with full capability and perfect memory simply computes a valid allocation. A central member of the monotonicity “coordination-required” class is thus not forcing under elimination. The other direction requires care, because here the dissociation is *representation-sensitive on the monotonicity side*. Separation-of-duties, quorum, independent confirmation, and blinding *can each be specified additively* — an attestation adds an approval, a vote aggregates, a replication adds a confirming result — and under that specification the monotonicity criterion rates them coordination-*free* while elimination rates them forcing. But they need not be written that way: a checking gate modeled as a *revisionary* approval that can reject prior work is non-monotone, and the near neighbor’s own reclassification explicitly labels such gates coordination-required. So under their revisionary representation these conditions *overlap* the non-monotone class. The point is therefore sharper than a clean dissociation: under elimination a checking gate is forcing because the approver is a distinct identity (C) or an error-decorrelated source (H) *regardless of how the specification is written*, whereas under monotonicity its classification flips with the representation. That representation-invariance is the real distinction, and it lives on the identity/interest/existence families the criterion cannot express in the first place.

Where the criteria co-extend — stated plainly. We do not claim clean orthogonality everywhere. The knowledge family is the exception: independent confirmation (H) and required ignorance (I) are epistemic, and a check that can *reject* is retractive, so on some cases “forcing under elimination” and “non-monotone under CALM” coincide. The distinguishing work is therefore carried by the identity, interest, and existence families (C, D, E, F, J), whose forcing requirements the monotonicity criterion does not pose — not by the epistemic family, which is where the two accounts are closest. We lead the “not a restatement” argument with the former for exactly that reason.

The account’s own stated scope, as corroboration. Consistent with all of the above, the near neighbor is explicit that it characterizes one dimension. It states that monotonicity is “one input to coordination need alongside epistemic, agent-level, incentive, and mechanism-fit considerations,” notes that epistemic interdependence “can create additional coordination needs that fall outside the scope of [the] result,” and, in its dependency mapping, that simultaneity “can matter for reasons other than non-monotonicity depending on the specification.” These map onto our families — knowledge (epistemic), interest/will and identity (incentive, agent-level), and spatiotemporal (simultaneity). We treat this as corroboration of a distinction already made on checkable grounds, not as its foundation: the positioning survives even with these statements removed, and we note the obvious risk that a future revision of that work could claim some of this territory directly. We position this work, then, not as a correction of the monotonicity criterion but as the *cross-family complement* to it.

A parity we concede. A defender of the monotonicity account will note, correctly, that *its*

classifications are representation-dependent (a label can change with how the decomposition is drawn) and that *ours* are too: conditions B and J depend on the readings of Section 6.1, and the count is eight-or-ten depending on grouping (Section 3.6). The parity is real; neither criterion is representation-free. But the *invariant cores* differ. No representation makes the monotonicity *criterion* condition on actor identity or post-destruction state: a quorum can be modeled in a monotone lattice — a grow-only set of votes — but the criterion then rates it coordination-*free*, blind to whether the votes come from one identity or k distinct ones, because monotonicity is about consistency under history growth, not about who acts. The distinct-identity requirement that makes quorum forcing under elimination is therefore invisible to the monotonicity classification however the task is drawn, whereas our representation-sensitivity is confined to two flagged conditions and never touches the identity or interest families. The orthogonality we claim is on that invariant core, and it survives the parity.

6.2.3 The common thread, and the difference

The prior camps and the near neighbor measure, in different vocabularies, the *information and fault structure* of a task: how outputs combine, when histories invalidate prior conclusions, what can be guaranteed under failure, or, in the planning account, which capabilities and state structure a single agent lacks. The elimination criterion does something different. It idealizes the single agent to perfection and asks what survives, and finds that what survives are conditions in which **being a single serial locus is itself the wall**: one body, one identity, one interest, one knowledge-state, one existence. That move — idealizing the single locus to perfection and reading off the structural residue across domains — is what no prior *cross-domain* account makes; its kinship with the group-productivity tradition (Section 6.1) is in the move, and its departure is the unbounded idealization and the reach across cryptography, law, aviation, and clinical method. That is what the taxonomy contributes.

7 Limitations

We state the limitations plainly, because the work’s credibility depends on the reader being able to tell what is established from what is offered.

- **The five-family axis is a synthesis, not an independent finding.** It is read off the reasons that recurred in grouping the conditions, not separately discovered. Treat it as an organizing model that fits the conditions well; the result is the conditions themselves, not the diagram.
- **The count is finer than the floor.** We report ten conditions, but a defensible conservative reading collapses the closest pairs to eight — or to seven, if the contested partitioned-secret condition (G) is also folded into required-ignorance (I) (Section 3.6). The within-family granularity (specifically the E/F and H/I distinctions, and the standing of the partitioned-secret condition G relative to required-ignorance I) is the genuine soft spot, and a skeptical reader is entitled to the coarser count.
- **Two conditions depend on a reading of the idealization.** Real-time deadline coverage (B) and independent survival (J) are forcing only under the readings in Section 6.1 — that unbounded patience does not extend the world’s clock, and that “never fails” does not confer invulnerability to external destruction. A stricter idealization that granted either would demote the corresponding condition. This is the principal interpretive choice in the work; we have argued the readings are the natural ones, but they are choices.

- **The thinnest conditions rest on the narrowest bases.** Required ignorance (I) and independent survival (J) are each supported by a smaller and less domain-diverse set of cases than the others. We judge them genuine and structurally clean, but they are the conditions most in need of further cases.
- **The taxonomy is scoped to safety-type forcing.** It characterizes when a single locus cannot produce a *correct* result at all. Availability/liveness under partition — a genuine, fault-tolerance-adjacent forcing in which the medium between loci, not the loci themselves, is the obstacle — is a deliberate scope boundary (Section 4), not a covered case. A unified account spanning the safety conditions here and a liveness/availability axis is the most significant piece of open work (Section 8).
- **The evidence is an illustrative survey, not an audited benchmark.** The cases were assembled with large-language-model assistance and without per-case human-expert validation; there is no released corpus or codebook. Every case carries a real, named referent and was subjected to the adversarial elimination test, but no domain expert independently certified each referent, and the second corpus is a sensitivity check (same criterion, same investigator), not an independent replication. The confidence is therefore in the *structure* — robust across the survey and the second pass (Section 5) — not in any individual case attribution, some of which may contain errors; the stronger tests of Section 5.3 (blind inter-rater agreement, a second operationalization, expert validation) are what a reproducible, mind-independence-grade study would add.

None of these undercuts the central claims — the taxonomy, the negative result, and the elimination criterion — but each bounds how far they should be pushed.

8 Discussion and implications

A diagnostic for multi-agent systems. The criterion yields a direct practical test for designers of multi-agent software. Before assigning a task to a team of agents, ask: *would a single agent with unbounded context, perfect memory, full reliability, and full capability still be unable to do this?* If the answer is no — if the only reasons to split it are speed, volume, or skill — then the task is in the degradation regime, and multiplicity is an optimization that scaling a single agent may erase. If the answer is yes, the task instantiates one of the conditions, and multiplicity is mandatory no matter how capable the individual agent becomes — which is exactly what “forcing” means. (For conditions B and J this presupposes the idealization readings of Section 6.1; a stricter idealization would move them below the band. And the test as stated screens for *correctness* forcing; availability under partition, Section 4, is a separate question it does not cover.) The conditions name *why*: the task needs two simultaneous effectors, two accountable identities, two opposed wills, two knowledge-states, or survival past one locus. This separates the coordination that scaling will dissolve from the coordination that no amount of scaling can.

The unifying idea. Across every domain examined — machinery and music, surgery and statute, markets and trials, cryptography and clinical method, aviation and inheritance — the forcing conditions reduce to a single shape: *a task forces more than one coordinating agent exactly when it requires two values of something a single locus of control can only ever be one of*. The diversity of the exemplars and the narrowness of the underlying shape are, together, what give the result the texture of something found rather than built (Section 5 and Section 7 say plainly what that texture rests on and what it does not yet establish).

Future work. Four directions follow naturally. First, and most important, a unified safety-and-liveness account: the present taxonomy characterizes safety-type forcing, and availability under partition (Section 4) is the salient forcing that lies outside it; an account joining the safety conditions here to a liveness/availability axis — with the CAP impossibility (Gilbert and Lynch) [19] supplying the latter’s teeth — would close the boundary the negative result exposes. Second, expert and field validation: certifying case attributions with domain experts, reporting blind inter-rater agreement on a held-out set (Section 5.3), and observing how the conditions manifest in real systems rather than curated cases. Third, a formal treatment of the boundary conditions B and J, making the idealization’s exogenous edges (Section 6.1) precise rather than argued. Fourth, a connection to common-knowledge impossibility: the knowledge family (partitioned secret, independent confirmation, required ignorance) sits close to results on what can be known and commonly known in a distributed setting [12], and a formal bridge there would sharpen the family most exposed to the charge of being merely epistemic restatement (Section 6.2).

9 Conclusion

We set out to separate the coordination a task *requires* from the coordination it merely *rewards*. The separation is achieved by elimination: posit a single agent freed of every endogenous limit — unbounded patience, unbounded durability, full capability, one serial locus — leave the world it acts in unchanged, and admit a condition only if that agent still cannot do the task while a coordinated set can. What survives is small and structured: a handful of conditions — eight, resolved more finely into ten — in five families, each a way that the *singleness* of the locus, not its weakness, is the wall. Redundancy for unreliable components — the first reason one reaches for — is shown to be degradation, not forcing, its genuine structure redistributing onto conditions the taxonomy already names, with availability under partition set aside as a separate boundary. The structure is stable across two separately assembled corpora: applying the same criterion a second time over a different corpus reproduced the negative result and the conditions up to resolution, and — once the two condition sets were reconciled — fell on the same family boundary. That second pass is a sensitivity check rather than independent replication, and Section 5.3 names the auditable tests a future study would add. As single agents grow more capable, the degradation reasons to coordinate will recede; these conditions are where coordination will remain — forced, not preferred.

References

- [1] J. D. Thompson. *Organizations in Action: Social Science Bases of Administrative Theory*. McGraw-Hill, 1967.
- [2] T. W. Malone and K. Crowston. The interdisciplinary study of coordination. *ACM Computing Surveys*, 26(1):87–119, 1994.
- [3] P. Puranam, M. Raveendran, and T. Knudsen. Organization design: The epistemic interdependence perspective. *Academy of Management Review*, 37(3):419–440, 2012.
- [4] H. Ju. When coordination is avoidable: A monotonicity analysis of organizational tasks. *arXiv:2602.18673*, 2026.
- [5] J. M. Hellerstein and P. Alvaro. Keeping CALM: When distributed consistency is easy. *Communications of the ACM*, 63(9):72–81, 2020.

- [6] J. M. Hellerstein. Complete CALM: A coordination criterion for specifications. *arXiv:2602.09435*, 2026.
- [7] P. Bailis, A. Fekete, M. J. Franklin, A. Ghodsi, J. M. Hellerstein, and I. Stoica. Coordination avoidance in database systems. *Proceedings of the VLDB Endowment*, 8(3):185–196, 2015.
- [8] M. Shapiro, N. Preguiça, C. Baquero, and M. Zawirski. Conflict-free replicated data types. In *Stabilization, Safety, and Security of Distributed Systems (SSS)*, 386–400, 2011.
- [9] M. J. Fischer, N. A. Lynch, and M. S. Paterson. Impossibility of distributed consensus with one faulty process. *Journal of the ACM*, 32(2):374–382, 1985.
- [10] L. Lamport, R. Shostak, and M. Pease. The Byzantine generals problem. *ACM TOPLAS*, 4(3):382–401, 1982.
- [11] M. Herlihy and N. Shavit. The topological structure of asynchronous computability. *Journal of the ACM*, 46(6):858–923, 1999.
- [12] J. Y. Halpern and Y. Moses. Knowledge and common knowledge in a distributed environment. *Journal of the ACM*, 37(3):549–587, 1990.
- [13] A. Shamir. How to share a secret. *Communications of the ACM*, 22(11):612–613, 1979.
- [14] D. D. Clark and D. R. Wilson. A comparison of commercial and military computer security policies. In *IEEE Symposium on Security and Privacy*, 184–195, 1987.
- [15] J. R. Searle. *The Construction of Social Reality*. Free Press, 1995.
- [16] M. Gilbert. *On Social Facts*. Routledge, 1989.
- [17] I. D. Steiner. *Group Process and Productivity*. Academic Press, 1972.
- [18] F. P. Brooks. *The Mythical Man-Month: Essays on Software Engineering*. Addison-Wesley, 1975.
- [19] S. Gilbert and N. A. Lynch. Brewer’s conjecture and the feasibility of consistent, available, partition-tolerant web services. *ACM SIGACT News*, 33(2):51–59, 2002. (Seth Gilbert — distinct from M. Gilbert, ref. 16.)
- [20] L. Lamport. The part-time parliament. *ACM Transactions on Computer Systems*, 16(2):133–169, 1998.
- [21] D. Ongaro and J. Ousterhout. In search of an understandable consensus algorithm. In *USENIX Annual Technical Conference (ATC)*, 305–319, 2014.
- [22] A. Avizienis. The N-version approach to fault-tolerant software. *IEEE Transactions on Software Engineering*, SE-11(12):1491–1501, 1985.
- [23] Y. Zhang, S. Sreedharan, and S. Kambhampati. A formal analysis of required cooperation in multi-agent planning. In *Proceedings of the International Conference on Automated Planning and Scheduling (ICAPS)*, 335–343, 2016. arXiv:1404.5643, 2014.
- [24] P. Jwalapuram et al. The illusion of multi-agent advantage. *arXiv:2606.13003*, 2026.