

What Is Situational Intelligence?

Why the category is needed, how it is defined, what distinguishes it from its neighbours – and what follows for anyone who builds a platform on it.

Harald Katzmair, PhD

Founder & Director, FAS Research, Vienna · harald.katzmair@fas.at · August 2026

ABSTRACT Organizations are richly equipped to pursue goals and poorly equipped to read the situations in which those goals must be pursued. This paper proposes *Situational Intelligence* as the missing category. It builds the definition from the ground up – what a situation is, what intelligence is, what follows for the capability and for a platform that instantiates it – and grounds the move in François Jullien's contrast between acting from a model and acting from the potential of a configuration. The category is specified by five commitments, a reference process of Map, Match, Mobilize and Learn, and seven jointly necessary criteria that make situational grounding inspectable. It treats the reading as a composite carried by four sources – own experience, stakeholder positions, artificial agents in defined roles, and models over evidence – and as something an interested adversary can attack directly. The paper marks the boundaries to adjacent constructs with falsifiable discriminant predictions, states how the claim could be measured and falsified, sets out what conformance requires of a platform, and closes with a self-assessment that any team can run on a decision it has already taken.

Keywords Situational Intelligence · situation · propensity · situational grounding · plurality · situated agency · agentic systems

Suggested citation Katzmair, H. (2026). What Is Situational Intelligence? FAS Research Working Paper 01/2026, Discussion Draft v5.2. Vienna: FAS Research.

Purpose This is a category description – not a market argument, and not a case study. It states what the category is, why it is needed, how it differs from its neighbours, how it could be measured and falsified, and what it requires of an implementation.

Status Proposed for independent adoption, critique, implementation and testing. The category should be adopted only if independent work establishes conceptual distinctiveness, reliable conformance testing and incremental value over simpler practices.

1 THE NEED: TWO WAYS OF ACTING

Two ideas of effective action have long competed. The first, dominant in European thought since the Greeks, begins with a model: form an ideal, fix the goal, apply it to the world, measure the gap. The second, which François Jullien reconstructs from classical Chinese strategic thinking, begins elsewhere. It starts from the configuration one is already in and from the potential latent in it – *shi*, the propensity of a situation to move in a particular direction before anyone plans anything. The strategist of the second kind does not impose a design. He evaluates the disposition of the field, detects where it is already inclined to go, and acts so that this inclination carries the action (Jullien, 1995, 2004).

The two are not rival theories of the world; they are two different competences. Modern organizations have industrialized the first one. Objectives, roadmaps, KPIs, business cases, execution tooling: an entire apparatus exists for setting goals and driving them through.

For the second competence, it would be wrong to claim that nothing exists. Situation awareness supplies perception and projection in fast, bounded environments (Endsley, 1995). Sensemaking explains how circumstances are turned into accounts that permit action (Weick, 1995). Deliberation platforms make the shape of disagreement visible at scale; facilitation practice holds difference in a room;

network analysis maps positions and dependencies; decision intelligence models and automates the decision itself. Each holds a genuine piece.

What is missing is not another piece. It is the joint – and specifically the joint between two things usually kept apart. On one side, formal analysis: network models, structured elicitation, scoring, forecasting, the apparatus of making a field calculable. On the other, the interpretive reading of a configuration: which way it leans, what it will bear, what a move here would cost there. Practitioners who are good at this do both at once, and the doing has no name. No established category requires that reading, option testing, mobilization and consequence-based updating be one connected, auditable capability – with a name, a standard, and a stated way of being wrong. Where these are joined today they are joined by a person: the senior figure with a feel for the field. That is what leaves when the person leaves, and what cannot be taught, audited, contracted for, or improved.

This gap has become expensive, for four reasons that increasingly appear together. Relevant knowledge is distributed across positions, so no single actor holds the picture. Interests, values and mandates differ, so there is no neutral vantage point. The field keeps moving while action is prepared – what Jullien calls silent transformation, change that completes itself before anyone names it (Jullien, 2011). And action changes the field, so every move al-

ters what remains possible for everyone else. Under these four conditions, more data, faster decisions and tighter execution do not correct a mistaken reading. They amplify it.

The same four conditions organize how strategies fail. We propose – as a taxonomy to be tested, not as an established empirical finding – that post-mortems sort into four recurring findings: **not seen** – the field was misread or read too narrowly; **not able** – capacity did not match the move; **not in time** – the opening closed while preparation ran; and **not wanted** – those who had to carry the action never committed to it. Each is the after-the-fact trace of something that could have been examined beforehand. What we propose is not procedural caution. It is these four failure modes turned into questions asked before acting rather than after.

We know how to pursue goals. Reading the situation in which a goal has to be pursued is a skill some people have and no institution owns – and what no institution owns cannot be taught, audited or improved.

2 FOUR DEFINITIONS

A category stands or falls on its terms. We build them in order: what a situation is, what intelligence is, what the two yield together, and what a platform must therefore do.

2.1 Situation

A situation is not an environment, a context, or a state. An environment surrounds an actor and is indifferent to it. A context is background. A state is a snapshot. A problem is already framed – and the framing is precisely what is at stake. A situation is none of these: it is a bounded configuration of actors and relations with a direction of movement, in which one actor's move alters what others can do.

DEFINITION 1 SITUATION

A situation is a temporally bounded and consequential configuration of actors, relationships, interests, knowledge, values, forces and constraints, organized around a focal question or possibility, in which one actor's moves alter the possibilities of others.

A situation is not given whole. Its boundary, focal concern, time horizon and relevant positions are analytical choices that must be declared and remain open to revision.

Three properties follow, and each does work later. A situation is *relational*: it consists of positions and dependencies, not of a list of factors (Figure 1). It is *dynamic*: it carries resources, momentum and propensities, together with tensions and gradients running across it (Lewin, 1951) – a still image of a situation is already a misreading. And it is *plural*: it looks different, and legitimately so, from the positions inside it. Actors carry different maps of the same field, those maps vary systematically by position, and accuracy about the structure is itself a source of power (Krackhardt, 1987, 1990).

2.2 Intelligence

The English word carries two lineages, and the category claims both. In *intellegere* – *inter* and *legere*, to read between, to pick out among – intelligence is the capacity to discern what lies between things rather than to inventory the things themselves. That is precisely what a relational field demands. In institutional usage – market intelligence, competitive intelligence – intelligence is a practice: the disciplined gathering and interpretation of what one must know in order to act, with sources, methods and standards. Situational Intelligence claims both senses: a capability of actors, and a practice that can be held to evidence.

DEFINITION 2 INTELLIGENCE, IN THIS CATEGORY

Intelligence is the capacity to read between – to discern relations, direction and possibility in a configuration – together with the practice of doing so from declared sources, by stated methods, to a standard that others can check.

It is carried by several kinds of reader, human and artificial; it operates on evidence that is usually unsharp; and it can be interfered with by actors who benefit from a misreading.

So understood, the intelligence in question is not exclusively cognitive, and it is not one faculty. Five capacities have to work together.

FIVE CAPACITIES OF THE READING

Discerning – reading between positions rather than inventorying factors: what depends on what, who can move whom, which relations bear load.

Unsharp classification – most of a live situation can be placed only approximately, between a lower and an upper bound, and must still be reasoned about rigorously (Pawlak, 1982). Waiting for sharpness is a way of arriving late.

Orientation – reading direction, not only structure: what the field is moving toward and away from, where the gradients run, how steep they are.

Anticipation – acting on a model of where the field is heading rather than a record of where it has been (Rosen, 1985).

Affective reading – treating interests, fears, trust and shame as data rather than noise (Salovey & Mayer, 1990), and recognizing configurations faster than they can be justified (Klein, 1998; Polanyi, 1966).

Orientation deserves emphasis, because it is what turns a map into a bearing. Lewin modelled a situation as a field of forces in which positive and negative valences produce movement toward and away (Lewin, 1951); Jullien's *shi* names the same phenomenon from the strategist's side, the disposition already inclining a configuration before anyone acts. The operative question is therefore never only *who stands where*, but *which way does this lean, how steeply, and how much of that lean can carry the move I have in mind*. A map without a bearing invites the most common error in strategy: a correct description of a field that has already begun moving somewhere else.

2.3 Whose intelligence?

In a consequential situation the capability is never held by one mind, and this is not a limitation to be engineered away. Four carriers contribute, each with privileged access and a characteristic failure mode; a reading that draws on only one of them is not economical but blind.

THE FOUR CARRIERS

My own experience – direct stake, history with the actors, tacit feel for what is possible here. *Fails through:* blind spots, interest in one's own framing, path dependence.

The stakeholders – positional knowledge no one else holds; each sees a part of the field visible only from where they stand. *Fails through:* partiality, exposure, what cannot safely be said.

Artificial agents in defined roles – tireless breadth, absent positions made testable, counter-positions sustained without fatigue. *Fails through:* fabrication, drift, and having nothing at stake.

Models over evidence – scale, recall, pattern across more cases than any participant has lived. *Fails through:* homogenization, training-set bias, and no access to the field itself.

The composition is the point. Human and artificial intelligence do not compete for the same job here: the carriers have different access, and their disagreements are diagnostic. Where the readings converge, the reading is located in shared discourse – which is not proof. Where they diverge, the divergence is the inquiry: unique field knowledge on one side, shared bias, model error or a missing position on the other.

2.4 Intelligence under adversarial conditions

Opposition in a situation operates on two levels, and only one of them is usually modelled. An actor may oppose the move: that is ordinary, it is a fact about the field, and testing for it is what Match does. An actor may also work against the *reading* – withholding what it knows, misdirecting attention, flooding the channel, capturing the framing before it is examined, or making it costly for particular positions to speak at all. The first is a condition to be navigated. The second is an attack on the capability itself.

This is not an exotic case. It is the ordinary condition of contested fields, and it is why the institutional lineage of the word has always carried a counterpart: intelligence practice implies counter-intelligence. A category built on reading must therefore treat the integrity of the reading as part of its object – who benefits from our misreading, which silences are produced rather than natural, which evidence arrives too conveniently. Situational grounding that cannot survive an interested adversary is not grounding; it is a rehearsal of what someone wanted us to conclude.

2.5 Situational Intelligence

DEFINITION 3 SITUATIONAL INTELLIGENCE

Situational Intelligence is the capability of an actor – or of actors acting in a shared situation – to read an evolving relational field, identify a viable strategic opening, mobilize appropriately governed action across difference, and update the reading from consequences.

Its object is **situational grounding**: an explicit basis in the best available evidence, a declared boundary and horizon, a relational account of likely support and resistance, legitimate authorization, and pre-declared signals for revision or stopping. Grounding is not certainty and not consensus. It is a contestable reason for acting here, now, in this configuration – and it decomposes into six components that the rest of the paper makes inspectable one by one: **framing, representation, orientation, viability, commitment** and **revision grounds**.

The capability is exercised in three modes. *Individual* Situational Intelligence is a single actor reading a field it does not control. *Collective* Situational Intelligence is several actors reading a shared field together. *Coordinated* Situational Intelligence couples the two: individual readings made mutually intelligible without being averaged. The two terms this paper leans on stand in a relation of scale. **Grounding** is a property of a particular action in a particular situation: present or absent this time. **Situated agency** is the standing capacity that repeated grounding produces – to act and adapt from within a particular field while remaining answerable to evidence, affected positions and consequences. Grounding is the unit; situated agency is what accumulates when the unit is achieved often enough to be relied upon.

2.6 A Situational Intelligence Platform

DEFINITION 4 SITUATIONAL INTELLIGENCE PLATFORM

A Situational Intelligence Platform is a governed software environment that instantiates the complete capability across repeated situations and makes its evidence, assumptions, mandates, actions and updates auditable.

It may use facilitation, structured data, network models, forecasting, generative AI or software agents. None of those technologies or methods alone is sufficient to define the category.

A platform is one form a conforming implementation can take, not the only one: the criteria apply equally to facilitated and mixed implementations. What a platform adds is repetition across situations and a record that does not depend on the memory of whoever was in the room.

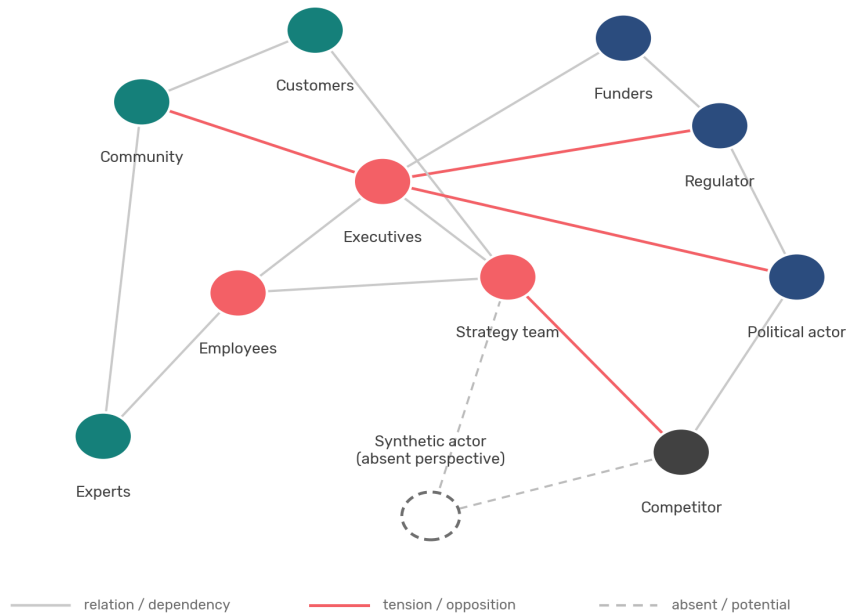


Figure 1. A situation as an analytically bounded relational field: heterogeneous positions, dependencies, tensions and relevant absences around a consequential focal question. The boundary is an analytical choice and must be declared.

3 WHAT THE CATEGORY CONSISTS OF

3.1 Five commitments

Relational: actors and possibilities are understood through positions and interdependencies. **Directional:** the field is read as something already in motion, with a lean that can be worked with, redirected or opposed – but not ignored. **Plural:** consequential differences remain visible rather than being averaged away. **Consequential:** action returns as evidence. **Contestable:** boundaries, sources, assumptions, mandates and revisions can be inspected and challenged.

The directional commitment is what carries Jullien's argument into the apparatus rather than leaving it in the preface. A reading that lists actors and relations without stating which way the configuration leans is a still image, and the paper's own claim is that a still image is already a misreading. Direction is therefore not an optional refinement of the map; it is a separate thing to be declared and to be worked about.

The plural commitment is the one most often mistaken for a courtesy, so it is worth stating why it is a necessity. Complex systems move through recurring phases of growth, conservation, release and reorganization – the adaptive cycle (Holling, 1986, 2001; Gunderson & Holling, 2002). Each phase rewards a different logic: entrepreneurial initiative in growth, procedural consolidation in conservation, precaution when limits assert themselves, solidarity and experiment in renewal. Cultural Theory's solidarities – hierarchy, individualism, egalitarianism and fatalism, with the hermit standing outside them (Thompson, Ellis & Wildavsky, 1990; Verweij & Thompson, 2006) – offer a plausible correspondence with these demands. We advance the mapping as a heuristic, not an established result; to our knowledge it has not been tested, and it is not tidy. Fatalism in particular has no phase in our mapping that rewards it,

which may be exactly why fatalistic readings are the first to be excluded from a room and the last to be disconfirmed. A field that silences a rationality is prepared for only one phase of its own history. Difference is retained not out of politeness but because the cycle will, sooner or later, visit every phase and demand the voice the previous phase preferred to suppress. Productively held dissonance supports innovation (Stark, 2009); cognitive diversity improves problem-solving under specified conditions (Page, 2007); and whether concerns surface at all depends on psychological safety (Edmondson, 1999).

3.2 The objection from Jullien

One objection presses from the direction this paper started. Jullien's polemic is against *modelization*: the prior model applied to the world. What follows here is a process, a set of criteria and a protocol – a model. We accept the bite, and answer with a distinction narrower than it may first appear. The criteria are minimum-*declaration* requirements, not minimum-*substance* ones. C2 requires that a boundary be drawn and declared; it does not say where to draw it. C3 requires that whatever lean the reading asserts be stated, given a strength, dated and made revisable; it does not say which way the field leans, and it is satisfied by the declaration we *cannot tell, and here is what would settle it*. What is modelled is the accountability of a reading – what must be visible, before the fact, for the reading to be contestable – rather than the field itself.

We do not claim this dissolves the objection. A declaration requirement still shapes what gets looked for, and an actor obliged to name a lean will find one. That is a real cost, and it is why the reflexive problem appears among the open questions in Section 5 rather than as something settled here. The concession should be drawn wider still: the *commitments* above – relational and directional in particular – are substance, not declaration. They assert in advance what

kind of thing a situation is. Jullien's objection retains its force there, and we have no answer to it beyond the record: fields read this way have proved easier to act in than fields read as lists.

3.3 The reference process

Four functions, run as a loop rather than a sequence (Figure 2). **Map** asks what is happening, how the field appears from relevant positions, what evidence supports those readings and what is missing; its output is a structured, contestable field model, not one shared truth. **Match** asks which move is desirable, legitimate and viable here – each option tested against capacity, timing, support and resistance, coalitions, authority, competing values, reversibility and unintended effects. A *strategic opening* is the result: not an opportunity that exists in an environment, but a move for which the required capacity, timing, relationships, legitimacy and room to act are already present or can credibly be assembled by this actor, here. An opportunity is a property of an environment; an opening is a property of the relation between an actor and a configuration. **Mobilize** converts the viable move into differentiated, authorized capacity: who leads, supports, legitimizes, monitors or retains the right to challenge, with mandates, dated checkpoints and stopping conditions. Alignment seeks agreement; mobilization makes action possible across difference. **Learn** resolves what was expected against what occurred, with ex-ante assumptions kept visible so that hindsight cannot rewrite the record.

3.4 Seven criteria

The seven criteria are jointly necessary for category conformance. They do not guarantee that an implementation is good – capability and impact require separate measurement – but each makes one component of grounding inspectable. Four concern the reading, two the action, one the return. Figure 3 sets out the whole architecture in one place – criteria, grounds, and the capability dimensions of Section 5 – including its one asymmetry: seven criteria open six components, because representation grounds is opened twice and from opposite sides.

C1 · Relevant plurality is elicited and preserved. Relevant positions, including absent, minority and opposing ones, are identified, and aggregation does not erase consequential difference.

→ *representation grounds: who is in the field model, and who is missing*

C2 · The situation is represented relationally and bounded explicitly. Actors, dependencies, influence, capacities and constraints are represented; focal question, boundary, horizon, sources and uncertainties are declared.

→ *framing grounds: what the situation is taken to be, and on what evidence*

C3 · The field's direction is read and declared. The reading states where the configuration is already inclined to move, how strongly, and how much of that movement the proposed action intends to use, redirect or oppose – or declares that the direction cannot be determined, and what

would settle it. The declared lean is revisited when the field moves.

→ *orientation grounds: which way this leans, how steeply, and how much of it we intend to borrow*

C4 · Positions become intelligible under governed visibility. Positions can be compared without requiring participants to adopt them, breaching agreed privacy, or treating perspective, evidence and authority as interchangeable.

→ *representation grounds, tested across positions: difference visible without being flattened*

C5 · Options are tested against the field ex ante. Support, resistance, capacity, legitimacy, harm, reversibility and failure conditions are examined; consequential expectations are recorded before outcomes are known.

→ *viability grounds: why this move is expected to work in this constellation*

C6 · Action is differentiated, authorized and bounded. Who acts, who authorizes, what may be delegated, where opposition remains, and what evidence escalates, redirects or stops execution – with named ownership and dated checkpoints.

→ *commitment grounds: who acts, on whose authority, within which bounds*

C7 · Action returns as evidence. How action altered relationships, perceptions and opportunities is recorded, and prior expectations are resolved against observable outcomes using pre-declared criteria and, for probabilistic forecasts, proper scoring rules (Brier, 1950). The changed field becomes the next input.

→ *revision grounds: what would count as being wrong, and what happens then*

Read the situation. Find the opening. Mobilize action. Learn from consequences.

4 WHAT IT IS NOT

The term has prior meanings. Operational analytics has used it for real-time fusion of assets, events and geospatial context (Stein, 2016); management literature uses *intelligence situationnelle* for context-sensitive managerial judgment (Tanneau et al., 2017). This proposal does not displace those uses. But a category earns its name by being distinguishable from the constructs closest to it, not merely from products. We therefore state the differences and attach to each a prediction that could show us wrong.

4.1 Situation awareness

Situation awareness is the nearest cousin and the likeliest confusion. It is defined for an operator in a bounded, fast environment where the state of the world is in principle knowable, and it proceeds through perception, comprehension and projection of that state (Endsley, 1995). Situational Intelligence differs on the object rather than the mechanism: the state is not merely unknown but *contested*, constituted by actors whose interests give them reason to represent it differently, and altered by the act of engaging it. Situation awareness improves as an operator's

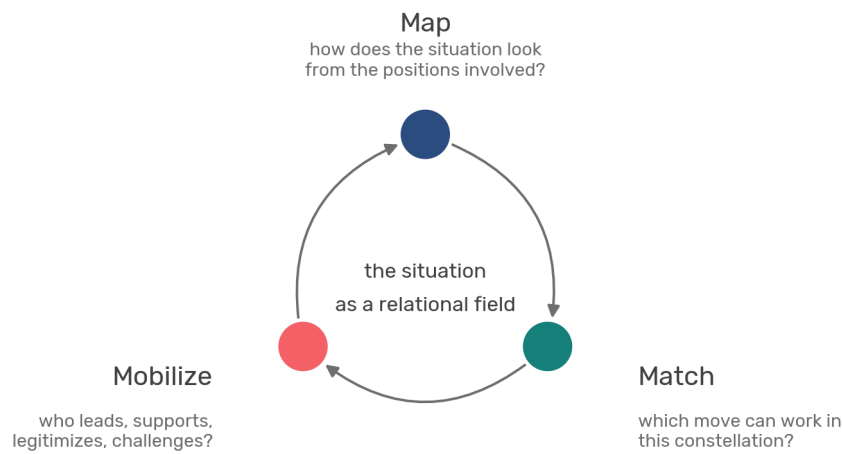


Figure 2. The reference process. Learning is the return path: action changes the field, and the changed field becomes the next input.

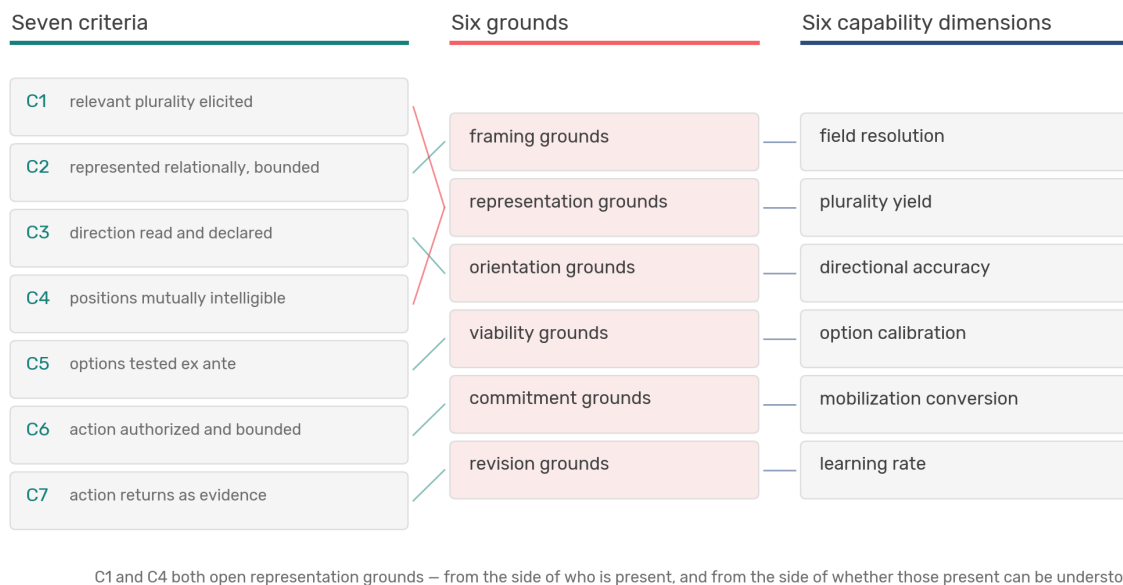


Figure 3. The architecture of the proposal: each criterion opens one component of situational grounding to challenge, and each component is evidenced by one capability dimension (Section 5).

model converges on the true state. Situational Intelligence has no single true state to converge on – only readings that can be made mutually intelligible, declared, and tested against what follows.

→ *discriminant prediction (double dissociation): interventions that raise situation-awareness scores – display design, workload reduction, alerting – should leave plurality yield and mobilization conversion unchanged; and interventions that raise plurality yield or mobilization conversion should leave situation-awareness scores unchanged. An intervention that moves both together would count against the distinction*

4.2 Sensemaking

Sensemaking is retrospective and explanatory: it accounts for how people bracket cues and build plausible narratives that permit action (Weick, 1995; Weick, Sutcliffe & Obstfeld, 2005). It describes something that happens whether or not anyone manages it. Sensemaking can certainly collapse – Weick's cosmology episode is precisely that – but it carries no conformance standard an implementation could be

held to and fail. Situational Intelligence is prospective and normative: it specifies what a reading must expose in order to justify action, and an implementation can be non-conforming.

→ *discriminant prediction: coded for sensemaking properties – bracketing, plausibility, enactment – conforming and non-conforming implementations should be indistinguishable, while grounding measures separate them cleanly. If sensemaking coding tracks conformance, the two constructs collapse into one*

4.3 Collective intelligence

Collective intelligence established a group-level factor that predicts performance across many tasks and is largely a property of the group itself (Woolley et al., 2010). Situational Intelligence is a property of the reading of one particular situation, and is expected to vary within the same group from field to field.

→ *discriminant prediction: a group's collective-intelligence factor should be comparatively stable – an assumption that itself needs testing, since temporal stability was not established in the*

original study – while its situational grounding varies with the field. If grounding turns out to be stable within groups and insensitive to the situation, it is collective intelligence renamed and this category should be withdrawn

4.4 Decision intelligence

Decision intelligence models and, increasingly, automates decisions, taking the organization as the decision-maker and treating variation among positions as noise to be resolved into a recommendation. That is a design choice, not a deficiency, and it is well suited to decisions an organization actually controls.

→ *discriminant prediction: decision-intelligence performance should hold steady whether or not authority over the decision is split, while grounding measures should fall sharply as authority fragments. If both degrade alike, the difference is one of degree, not of kind*

4.5 The category map

The same distinctions sort the practice classes that exist today, along two descriptive questions rather than two claims of merit (Figure 4): what is the unit of analysis, and what happens to difference. Both are answerable about a class of practice without asking whether that practice is any good, and each placement is one its occupants should be willing to state of themselves. Prediction markets are the limiting case in one corner – exquisite calibration purchased by compressing every perspective into a single price. Deliberation platforms preserve difference genuinely but take a population's opinion space, not a situation with authority in it, as their object. Facilitation practice works on the right unit and does preserve difference; what it does not leave behind is a record independent of the practitioner who held the room, which is why it appears here as a practice rather than as an instrument.

NON-INSTANCES

A dashboard that fuses real-time data but represents no plural positions and tests no strategic move.

A workshop that creates insight or alignment but records no governed action and no consequence-based update.

An agent platform that plans and executes without a contestable situation model, mandate and stopping logic.

A stakeholder map or scenario exercise that ends before owned action and longitudinal learning.

4.6 Postures, and Boyd

Three further vocabularies name postures rather than capabilities. **Readiness** is ex ante, **responsiveness** in tempore, **resilience** ex post – of which only the last has an established technical sense, the capacity of a system to absorb disturbance and persist (Holling, 1973); the ordering into a triad is our own. All three are states a system holds to some degree; none says how the state is produced. We propose Situational Intelligence as one capability behind all three: a field well read confers readiness, movement detected in time confers responsiveness, and plurality preserved for the next phase of the cycle confers resilience.

Boyd's OODA loop deserves a more careful treatment than it usually receives, including from us. The popular reading – a fast decision cycle in a cockpit – understates Boyd considerably; Osinga (2007) shows that the Orientation node already carries culture, heritage and prior experience, and that *Organic Design for Command and Control* is explicitly organizational, built on shared implicit orientation across many actors. Our difference with Boyd is therefore neither tempo nor scope. It is purpose. Boyd's shared orientation is an asset to be converged, so that an organization can act faster than an adversary can respond. Here, consequential divergence between positions is retained rather than converged, because it is diagnostic evidence about the field; and legitimacy, authorization and the standing of affected positions are part of the object rather than constraints upon it.

5 HOW THE CLAIM CAN BE TESTED

A proposed category earns standing through independent validation, not repetition. Measurement must separate three questions that are routinely conflated. **Conformance:** does an implementation belong to the category – are all seven criteria present and auditable? **Capability:** how well does it perform the constituent functions? **Impact:** does using it improve outcomes against a credible alternative? Only the second is addressed by the six dimensions below, and each evidences one component of grounding.

SIX CAPABILITY DIMENSIONS

Field resolution – does the map hold against independent evidence? (*framing*)

Directional accuracy – did the field move where the reading said it leaned, and was the declared lean revised when it did not? (*orientation*)

Plurality yield – do protected contributions and examined divergences materially change map, options, authorization or action? (*representation*)

Option calibration – do ex-ante expectations resolve well against reference classes and proper scoring rules? (*viability*)

Mobilization conversion – what share of named commitments and intended movements is realized by the horizon? (*commitment*)

Learning rate – how quickly does disconfirming evidence update the reading, and do recurrent blind spots decline? (*revision*)

A composite score would be premature. An open protocol should first publish operational definitions, coding rules, reference classes, baselines, uncertainty and a change log; validation should include independent implementations, inter-rater reliability, multimethod construct tests and comparison with strong alternatives (Cronbach & Meehl, 1955; Campbell & Fiske, 1959).

Six problems remain genuinely open, and we would rather name them than survive by vagueness. *Discriminant validity:* can the construct be told apart empirically from situation awareness (Endsley, 1995), sensemaking (Weick, 1995), collective intelligence (Woolley et al., 2010) and decision in-



Figure 4. Practice classes along two descriptive questions. Neither axis ranks quality; each placement is one its occupants should be willing to state of themselves. The upper-right region is unoccupied not because it is superior but because it demands a combination each neighbouring class has its own good reasons to avoid.

telligence? *Boundary reliability:* can independent analysts delimit the same focal situation, positions and horizon, and which exclusions persist systematically? *Plurality, power and safety:* when does preserving disagreement improve outcomes, and when does it produce false equivalence, exposure, delay or capture by the powerful? *Performativity:* how can forecast skill be distinguished from intervention effect when acting on a reading changes the outcome being predicted? *Adversarial robustness:* how much degradation of the reading – strategic silence, flooding, framing capture – can a conforming process absorb before its grounding claim fails, and can that threshold be measured rather than asserted? And sixth, pressing hardest on Section 3, *the reflexive problem:* does a declaration requirement in fact leave the reading free, or does it quietly prescribe what gets looked for – an actor obliged to name a lean will name one – so that conformance reintroduces the very modelization the category was built to avoid?

5.1 What we can and cannot claim today

A word on our own evidence, since the paper asks others to be falsifiable. FAS Research has worked on situations of this kind for three decades and operates a platform built to these criteria. We have practice; we do not yet have published validation. No conformance assessment has been independently replicated. No capability dimension has pub-

lished inter-rater reliability. No impact study against a credible alternative has been completed. Every performance claim implied anywhere in this paper should therefore be read as a hypothesis about our own work rather than as evidence for it. The protocol and the first results, including null findings, will be published separately, and the category deserves to be judged on those rather than on this paper.

WHAT WOULD COUNT AGAINST THE CATEGORY

If the construct cannot be reliably distinguished from established adjacent capabilities, it should not be treated as a separate category.

If the integrated process adds no value over simpler practices under fair comparison, it should remain a descriptive bundle, not a performance claim.

If independent raters cannot apply the conformance criteria consistently, the boundary must be revised before market adoption.

6 WHAT FOLLOWS FOR A PLATFORM

If the category holds, it constrains what may be built under its name. Five consequences follow directly.

Conformance is about the joint, not the components. A platform conforms when all seven criteria are instantiated

in one connected, auditable process. This is not a disqualification of component suppliers: a network map, a deliberation tool, a forecasting module and an agent runtime are inputs the category depends on, and most conforming implementations will be assembled from them. What conformance adds is the requirement that the joins exist and can be inspected. We would rather be measured on this than assert it: FAS Research undertakes to publish conformance assessments of rival implementations alongside its own, including any that score better.

Plurality has to be engineered, not hoped for. Preserving difference is a mechanism, not an attitude: protected elicitation, origin kept visible on every contribution, results that show the spread before any average, and a way of holding a split open when it runs between the people in the room rather than merely across the data.

AI is an enabler, never the grounds. Artificial intelligence lowers the cost and raises the frequency of mapping, comparison, simulation and monitoring, but it does not confer grounding. Human-AI combinations are not automatically superior: a meta-analysis of 106 experiments found average performance below the better of the human-only and AI-only baselines ($g = -0.23$), with losses concentrated in decision tasks ($g = -0.27$) and gains in open-response creation tasks ($g = +0.19$); combinations did improve on humans alone ($g = 0.64$) (Vaccaro, Almaatouq & Malone, 2024). The distinction bears on the design proposed here: the tasks assigned to machines – structuring, comparing, generating counter-positions – are creation-type on that taxonomy, and the tasks withheld are decisions. That is a design intention, not a result; whether an implementation's task assignment actually falls on the favourable side of the line is something conformance testing should check rather than assume. Language models can simulate recognizable responses with documented but bounded fidelity – interview-grounded agents reproduced held-out survey responses at 83% of participants' own two-week test-retest consistency, 86% when surveys were added, against 74% for demographics-only agents (Park et al., 2024) – which supports structured rehearsal, not synthetic authority. A stakeholder simulation is a testable hypothesis about a possible response, never the stakeholder's voice. Two risks deserve standing attention: *homogenization*, when people and agents converge on the same discourse – a risk that grows precisely because models can reproduce fine-grained positions convincingly (Argyle et al., 2023; Bail, 2024), and *performativity*, when action changes the situation so that later outcomes are no longer independent test data. Deploying an agent is itself a strategic move in the field, and must be tested as one.

Governance is constitutive, not compliance. Because the category processes speech, positions, stakeholder models, decisions and mandates, it requires data minimization, purpose limitation, informed participation, role-based access, explicit retention and documented provenance. Deliberation data must not be repurposed for covert people scoring, performance monitoring or surveillance. Agents may make only bounded decisions whose mandate, authorization, escalation, stopping conditions and accountability are inspectable, alongside – not instead of – established AI management and oversight controls (ISO/IEC 42001:2023; NIST, 2026).

The integrity of the reading is part of the product. If opposition can operate on the reading and not only on the move, a platform must make the conditions of its own evidence inspectable: who contributed and who did not, which positions were invited and which declined, where contributions were protected and where attribution was visible, and what arrived from a model rather than from the field. Origin kept visible on every item is not bookkeeping. It is what allows a later reader to ask whether the grounding was earned or manufactured.

Conforming implementations may be human-facilitated, software-mediated, AI-assisted or agentic, and may run as deep sessions for contested choices, as light recurring rituals, or as continuous processes with longitudinal updates. No session format, product architecture or vendor defines the category. External validity begins when independent organizations implement the same criteria, fail the same tests, and report against the same open protocol.

Fields keep moving whether or not anyone reads them. Situational Intelligence is the discipline of acting inside that movement – on solid grounds, with others, and with the willingness to be corrected by what happens next.

Disclosure. Harald Katzmair, PhD, is Founder & Director of FAS Research, which develops and operates SituationR®, a commercial implementation of practices discussed here. SituationR is not the category, and no product feature is a category requirement unless the same requirement can be applied to every implementation. Product evidence is not treated as validation of the category. FAS Research commits to publishing definitions, protocols, revisions and empirical results including null findings, and invites independent replication and rival implementations; stewardship should become institutionally independent of any vendor once a credible multi-implementer community exists. The category name and criteria are offered for open use and criticism.

Correspondence. Harald Katzmair · FAS Research, Vienna · harald.katzmair@fas.at · fas-research.ai

7 APPENDIX: A SELF-ASSESSMENT

The criteria are of little use as an argument if they cannot be applied. What follows is the shortest form we know of. Take a consequential decision your organization has *already* taken, ideally six to twelve months ago, so that consequences exist. Answer seven questions from the record – not from memory. Score each: 2 if the record shows it and an independent reader could reconstruct the reasoning and name what would have counted as being wrong; 1 if it happened but was not recorded, or was recorded too thinly to reconstruct; 0 if neither. Mark X where something was recorded and later turned out to be wrong – X is not a penalty; it is the only mark that distinguishes a reading from a filing system. A profile carrying no X marks twelve months on means one of two things, and both are worth knowing: nothing was predicted, or nothing was checked.

Two limits should be stated plainly. The instrument measures the *auditability* of a reading, not its quality: a well-documented catastrophic misreading scores highly, which

is why the X marks matter more than the total. And the record of a contested decision is written by whoever prevailed, so question 4 in particular cannot be scored from documents at all – it has to be asked of the parties, live, and a score obtained any other way should be treated as missing. The instrument has no published reliability; a score is a conversation, not a certification.

THE SEVEN QUESTIONS (SI-7)

1 · Plurality. Which positions were consulted, which were knowingly absent, and what was done about the absences? (C1)

2 · Framing. Where was the boundary of the situation drawn, by whom, on what evidence, and what was left outside it? (C2)

3 · Direction. What did we say the field was already moving toward, how strongly, and how much of that movement were we counting on? (C3)

4 · Intelligibility. Could each party state the others' position in a form those others would accept? (C4)

5 · Options ex ante. What did we write down, before the outcome was known, about expected support, resistance and failure conditions? (C5)

6 · Authorized action. Who owned each move, under whose mandate, with which checkpoint, and what would have stopped it? (C6)

7 · Return. Which expectations have since been resolved against what actually happened, and what changed in our reading as a result? (C7)

The exercise takes about ninety minutes, and the interesting result is never the total but the shape of the profile. Our own unsystematic impression from unpublished use – offered as an expectation to be tested, not as a finding – is that scores cluster high on question 6 and low on 3, 5 and 7: organizations can say who was responsible, and cannot say what they expected, which way they thought the field was leaning, or what they have learned since. If that pattern holds under systematic use it would matter, because a reading that cannot be scored on questions 3, 5 and 7 was not grounded but asserted – and the record, by itself, cannot tell the difference.

REFERENCES

- Argyle, L. P., Busby, E. C., Fulda, N., Gubler, J. R., Rytting, C., & Wingate, D. (2023). Out of one, many: Using language models to simulate human samples. *Political Analysis*, 31(3), 337–351.
- Bail, C. A. (2024). Can Generative AI improve social science? *Proceedings of the National Academy of Sciences*, 121(21).
- Brier, G. W. (1950). Verification of forecasts expressed in terms of probability. *Monthly Weather Review*, 78(1), 1–3.
- Campbell, D. T., & Fiske, D. W. (1959). Convergent and discriminant validation by the multitrait-multimethod matrix. *Psychological Bulletin*, 56(2), 81–105.
- Cronbach, L. J., & Meehl, P. E. (1955). Construct validity in psychological tests. *Psychological Bulletin*, 52(4), 281–302.
- Edmondson, A. (1999). Psychological safety and learning behavior in work teams. *Administrative Science Quarterly*, 44(2), 350–383.
- Endsley, M. R. (1995). Toward a theory of situation awareness in dynamic systems. *Human Factors*, 37(1), 32–64.
- Gunderson, L. H., & Holling, C. S. (Eds.) (2002). *Panarchy: Understanding Transformations in Human and Natural Systems*. Island Press.
- Holling, C. S. (1973). Resilience and stability of ecological systems. *Annual Review of Ecology and Systematics*, 4, 1–23.
- Holling, C. S. (1986). The resilience of terrestrial ecosystems: Local surprise and global change. In W. C. Clark & R. E. Munn (Eds.), *Sustainable Development of the Biosphere* (pp. 292–320). Cambridge University Press.
- Holling, C. S. (2001). Understanding the complexity of economic, ecological, and social systems. *Ecosystems*, 4(5), 390–405.
- ISO/IEC. (2023). *ISO/IEC 42001:2023, Information technology – Artificial intelligence – Management system*.
- Jullien, F. (1995). *The Propensity of Things: Toward a History of Efficacy in China*. Zone Books.
- Jullien, F. (2004). *A Treatise on Efficacy: Between Western and Chinese Thinking*. University of Hawai'i Press.
- Jullien, F. (2011). *The Silent Transformations*. Seagull Books.
- Klein, G. (1998). *Sources of Power: How People Make Decisions*. MIT Press.
- Krackhardt, D. (1987). Cognitive social structures. *Social Networks*, 9(2), 109–134.
- Krackhardt, D. (1990). Assessing the political landscape: Structure, cognition, and power in organizations. *Administrative Science Quarterly*, 35(2), 342–369.
- Lewin, K. (1951). *Field Theory in Social Science: Selected Theoretical Papers*. Harper.
- National Institute of Standards and Technology. (2026). *Accelerating the Adoption of Software and Artificial Intelligence Agent Identity and Authorization: Concept Paper*. February 2026.
- Osinga, F. P. B. (2007). *Science, Strategy and War: The Strategic Theory of John Boyd*. Routledge.
- Page, S. E. (2007). *The Difference: How the Power of Diversity Creates Better Groups, Firms, Schools, and Societies*. Princeton University Press.
- Park, J. S., Zou, C. Q., Shaw, A., Hill, B. M., Cai, C., Price, M. R., Liang, P., Willer, R., & Bernstein, M. S. (2024). Generative agent simulations of 1,000 people. arXiv:2411.10109 (figures cited from the current version).
- Pawlak, Z. (1982). Rough sets. *International Journal of Computer and Information Sciences*, 11(5), 341–356.
- Polanyi, M. (1966). *The Tacit Dimension*. Doubleday.
- Rosen, R. (1985). *Anticipatory Systems: Philosophical, Mathematical and Methodological Foundations*. Pergamon Press.
- Salovey, P., & Mayer, J. D. (1990). Emotional intelligence. *Imagination, Cognition and Personality*, 9(3), 185–211.
- Stark, D. (2009). *The Sense of Dissonance: Accounts of Worth in Economic Life*. Princeton University Press.
- Stein, R. (2016). Making effective decisions in real time: Situational intelligence brings together advanced analytics, data visualization and IoT. *Analytics Magazine*, 5 September 2016.
- Tanneau, C., Delahaie, P., Bonnefous, A.-M., & Fiol, M. (2017). *L'intelligence situationnelle: 50 situations de management décryptées, 67 fiches concepts*. Eyrolles.
- Thompson, M., Ellis, R., & Wildavsky, A. (1990). *Cultural Theory*. Westview Press.
- Vaccaro, M., Almaatouq, A., & Malone, T. W. (2024). When combinations of humans and AI are useful: A systematic review and meta-analysis. *Nature Human Behaviour*, 8, 2293–2303.
- Verweij, M., & Thompson, M. (Eds.) (2006). *Clumsy Solutions for a Complex World: Governance, Politics and Plural Perceptions*. Palgrave Macmillan.
- Weick, K. E. (1995). *Sensemaking in Organizations*. Sage.
- Weick, K. E., Sutcliffe, K. M., & Obstfeld, D. (2005). Organizing and the process of sensemaking. *Organization Science*, 16(4), 409–421.
- Woolley, A. W., Chabris, C. F., Pentland, A., Hashmi, N., & Malone, T. W. (2010). Evidence for a collective intelligence factor in the performance of human groups. *Science*, 330(6004), 686–688.