



FROM THE RESEARCH BENCH

Hybrid intelligence needs a boundary between capability and outcome

Hybrid intelligence names complementary human and machine capability, not proof of better performance: define the construct before measuring the outcome.

Sinan Isoglu, MBA (Quantic)

Commercial growth leader, lecturer and doctoral researcher

PUBLISHED

1 September 2026

READING TIME

7 min

WORDS

1,633

<https://isoglu.com/insights/journal/hybrid-intelligence-needs-a-boundary-between-capability-and-outcome/>

Management summary	2
Why is hybrid intelligence a descriptive label rather than a validated scientific construct?	3
What foundational definitional boundaries did Dellermann et al. establish for hybrid systems?	5
Why must system capability be strictly separated from realized commercial performance?	5
Why is ongoing operational governance an intrinsic component of hybrid capability claims?	5
How does true symbiotic hybrid intelligence differ from static task decomposition?	6
How should commercial organizations structure a synthetic hybrid intelligence evaluation?	6
Which three conventional claims regarding AI superiority must leadership resist?	7
Where are the empirical boundaries of hybrid intelligence theory?	7
References	8

MANAGEMENT SUMMARY

Hybrid intelligence is easy to use as a flattering label for any team that uses AI. Dellermann, Ebel, Söllner, and Leimeister define it more narrowly as a division of labor that combines complementary human and artificial strengths. They distinguish heterogeneous human-machine agents from homogeneous collective intelligence and identify trust, transparency, governance, and incentives as open issues for complex and uncertain work. This article separates the construct, its observable evidence, and the team outcome that would have to be measured. It turns the distinction into a synthetic evaluation card, distinct from task-allocation advice, and does not claim that a current team outperforms either humans or machines alone.

Keywords: Hybrid intelligence · Collective intelligence · Human-machine collaboration · AI evaluation

Hybrid intelligence is not proof that a team performed better.

The short answer is that hybrid intelligence names a proposed division of complementary human and artificial capability, while performance is a separate claim that needs an observed outcome. A team can use an AI system without arranging complementary work. It can arrange complementary work without measuring whether the team decision improved.

Dellermann, Ebel, Söllner, and Leimeister define hybrid intelligence as a division of labor that uses complementary strengths of human and artificial intelligence. They distinguish heterogeneous human-machine agents from homogeneous collective intelligence and identify trust, transparency, governance, and incentives as open design issues for complex and uncertain tasks.

Jarrahi (2018) describes the complementary division in organizational decision making more specifically: machines can contribute analytical processing and pattern recognition, while humans contribute contextual understanding, intuition, judgement, and responsibility. Seeber et al. (2020) extend the boundary to AI teammates, where task allocation, coordination, communication, trust, accountability, and team outcomes become design questions. Together, these sources support a construct boundary without turning complementarity into a performance result.

That distinction is useful because AI language often collapses three different statements:

- 1 the system has a technical capability;
- 2 the human and machine contributions are complementary;
- 3 the combined team produces a better outcome.

Only the first may be visible in a product description. The second needs a work design. The third needs a comparison and an outcome.

Why is hybrid intelligence a descriptive label rather than a validated scientific construct?

Calling a group a hybrid-intelligence system does not explain which agents are involved, what each contributes, or how their outputs interact. A construct becomes reviewable when its elements can be identified and an evaluator can say what would count as evidence.

Table 1 Why is hybrid intelligence a descriptive label rather than a validated scientific construct?

Construct element	Question	Evidence still needed
Agent heterogeneity	Are human and artificial agents contributing different capabilities?	A description of the actual contributions
Complementarity	Does each agent provide something the other does not provide as well in this task?	A task-level comparison
Division of labor	Where is work assigned, combined, or handed off?	A visible workflow
Governance	How are trust, transparency, incentives, and responsibility handled?	Rules and review records
Team outcome	What should improve because the agents work together?	A declared outcome and comparison

Table from this essay. Sources and interpretation are given in the article.

The last row must not be inferred from the first four. A clear division of labor can still produce a poor decision. A strong team outcome can arise for reasons unrelated to the AI contribution. The construct boundary prevents a performance claim from hiding behind an attractive label.

Figure 1 The hybrid-intelligence construct boundary

Construct question	Observable evidence	Missing evidence	Outcome measure	Interpretation limit
Who contributes?	Human and machine roles are described	Contribution quality is unknown	None yet	A role list is not complementarity
What is complementary?	Capabilities differ for the task	No matched task comparison	Task-level performance	Difference is not superiority
How is work governed?	Trust, transparency, and responsibility are named	Incentives and challenge paths are unclear	Review and correction rate	Governance is not proof of safety
What improved?	A team outcome is declared	No baseline or comparison	Declared decision outcome	A label cannot supply the result

Author's synthetic evaluation framework grounded in Dellermann, Ebel, Söllner and Leimeister (2019), Jarrahi (2018), and Seeber et al. (2020). The rows are illustrative and do not score a current team.

The card does not measure hybrid intelligence. It records the missing bridge between a concept and an outcome. A team may fill the first three rows and still have no evidence in the fourth.

What foundational definitional boundaries did Dellermann et al. establish for hybrid systems?

Dellermann and colleagues describe hybrid intelligence as a division of labor that uses complementary strengths of human and artificial intelligence. The definition matters because it does not say that any human-machine contact is hybrid intelligence. It requires a relationship between capabilities and a task.

They distinguish hybrid intelligence from collective intelligence among homogeneous human or animal groups. The distinction is not a judgement that one arrangement is better. It points to different design conditions. Heterogeneous agents may have different information, speed, error patterns, incentives, and ways of communicating. Those differences create a coordination problem that a label alone does not solve.

For complex and uncertain tasks, the source places human creativity, empathy, intuition, and domain knowledge alongside machine analytic capacity. Jarrahi (2018) similarly separates machine pattern recognition from human contextual judgement and responsibility. Seeber et al. (2020) show why this division also creates team questions about communication, trust, accountability, and outcomes. Those issues belong in the construct boundary because a technically capable system can be difficult to challenge, explain, or assign responsibility around.

Why must system capability be strictly separated from realized commercial performance?

Consider a synthetic review process in which a machine searches a bounded document set and a person interprets the result. The system may have a useful search capability. The human may add contextual judgement. That is evidence of a designed division of contributions, not evidence that the combined decision is better.

The team outcome might be faster review, fewer missed items, better explanation, or a different decision quality measure. Each outcome requires a declared comparison. Faster is not automatically better. Fewer missed items may come with more false alarms. Better explanation can trade off against time. The evaluation must name what the team is trying to improve and what it will sacrifice or constrain.

This is why “hybrid intelligence outperforms humans and machines” is too strong as a general sentence. The source provides a construct and a design problem. It does not provide a universal performance result for every task, team, or system.

Why is ongoing operational governance an intrinsic component of hybrid capability claims?

Trust and transparency are sometimes added after the technical capability has been described. The source places them inside the open design problem. That placement changes the evaluation question.

If a human cannot understand what the machine considered, challenge its output, or identify who owns the final decision, the team may have complementary capabilities but a weak governance arrangement. If incentives reward agreement with the system, the observed team outcome may reflect compliance rather than useful complementarity.

Use four governance prompts:

Table 3 Why is ongoing operational governance an intrinsic component of hybrid capability claims?

Prompt	Why it belongs in the construct boundary
What can the person inspect?	Transparency determines whether the machine contribution is challengeable
What can the person override?	Control determines whether human judgement remains meaningful
Who owns the decision?	Accountability prevents the machine label from becoming a responsibility gap
What is rewarded?	Incentives can change whether disagreement and correction are possible

Table from this essay. Sources and interpretation are given in the article.

These are author-owned prompts grounded in the source’s open issues. They do not prove that a governed system is safe or effective.

How does true symbiotic hybrid intelligence differ from static task decomposition?

An implementation plan may ask which tasks belong to a human or a machine, where the handoff occurs, and how coordination works. That is a task-allocation problem. This article asks a prior question: does the arrangement satisfy the construct boundary for hybrid intelligence, and what evidence would be needed before claiming a team outcome?

The distinction matters in a large program. A task map can be clear while the capability relationship is not complementary. A construct can be well defined while the handoff is unusable. A team can report an outcome while the baseline is missing. Each is a different gap.

How should commercial organizations structure a synthetic hybrid intelligence evaluation?

Before using the label in a strategy or research document, record:

- 1 the human and artificial agents involved;
- 2 the task for which their capabilities are being compared;
- 3 the complementary contribution of each agent;
- 4 the work boundary and information exchanged;
- 5 the trust, transparency, governance, and incentive conditions;
- 6 the team outcome and its baseline or comparison;
- 7 the failure and correction path;
- 8 the evidence that the outcome belongs to the combined arrangement.

The last field is the hardest. A team outcome can improve because a new process, training program, or selection rule changed at the same time. The label should not receive credit for every simultaneous change.

Which three conventional claims regarding AI superiority must leadership resist?

First, do not say “we use AI, therefore we have hybrid intelligence.” Use the construct definition.

Second, do not say “hybrid intelligence improves performance” without a task-level outcome and comparison. Capability complementarity is not a result.

Third, do not say “governance is solved” because trust, transparency, and accountability are named. They are design questions that need evidence.

For adjacent decisions, compare what AI changes in revenue operations with the playbook evidence boundary.

Where are the empirical boundaries of hybrid intelligence theory?

Dellermann and colleagues define hybrid intelligence as a division of labor using complementary human and artificial strengths, distinguish heterogeneous agents from homogeneous collective intelligence, and identify governance issues for complex and uncertain work. Jarrahi provides a human-AI division-of-labor boundary, while Seeber et al. provide a team-collaboration boundary. The three sources do not evaluate a current team or guarantee superior performance. The construct card is an author-owned translation. It keeps capability, governance, and outcome separate so that a useful design does not become an unsupported performance claim.

REFERENCES

- Dellermann, D., Ebel, P., Söllner, M., & Leimeister, J. M. (2019). Hybrid intelligence. *Business & Information Systems Engineering*, 61, 637-643. DOI
- Jarrah, M. H. (2018). Artificial intelligence and the future of work: Human-AI symbiosis in organizational decision making. *Business Horizons*, 61(4), 577-586. DOI
- Seeber, I., Bittner, E., Briggs, R. O., de Vreede, T., de Vreede, G.-J., Elkins, A., Maier, R., Merz, A. B., Oestere-Reiss, S., Randrup, N., & Reuter, M. (2020). Machines as teammates: A research agenda on AI in team collaboration. *Information & Management*, 57(2), 103174. DOI

END OF ESSAY

HOW TO CITE THIS ESSAY

Isoglu, S. (2026, September 1). Hybrid intelligence needs a boundary between capability and outcome. Sinan Isoglu.
<https://isoglu.com/insights/journal/hybrid-intelligence-needs-a-boundary-between-capability-and-outcome/>

KEEP READING

From the research bench

Human and machine should divide the work

<https://isoglu.com/insights/journal/human-and-machine-should-divide-the-work/>

From the research bench

What is external validity? A result needs a transfer boundary

<https://isoglu.com/insights/journal/what-is-external-validity/>

From the research bench

What is interference? When one unit changes another unit's outcome

<https://isoglu.com/insights/journal/what-is-interference/>



Sinan Isoglu, MBA (Quantic)
Commercial growth leader, lecturer and doctoral researcher