

PROJECT INTENT WORKING PAPER SERIES

The Clarity Mandate

First Edition

Alan Florschuetz

projectintent.com

Recognition

Everyone who has worked effectively with AI has noticed something familiar.

You provide clear intent. You define constraints. You let the system execute. You evaluate the output. You correct course.

It works remarkably well — when you do it that way.

Now think about the best leader you ever worked for.

They told you what mattered and why. They were clear about the boundaries. They trusted you to figure out how. They evaluated results, not process. They corrected quickly when things drifted.

The pattern is the same.

Not metaphorically. Structurally.

The Pattern

Both effective leadership and effective AI collaboration follow a single operating sequence:

1. **Intent** — define the desired outcome and why it matters.
2. **Constraints** — establish the boundaries that cannot be crossed.
3. **Autonomous execution** — trust the executor to discover the path.
4. **Evidence** — evaluate what actually happened.
5. **Correction** — adjust direction based on reality.

This is not a new model. It is the oldest model that works. Military doctrine calls it mission command. Management science calls it management by objectives. Agile attempted it with sprint goals. The pattern recurs because it reflects something fundamental about how capable agents — human or artificial — produce good outcomes under uncertainty.

What changes with AI is not the model.

What changes is that AI makes the model visible.

Why It Was Hidden

For decades, most organizations operated with a different pattern:

1. Strategy (somewhere at the top)
2. Detailed instructions (cascaded downward)
3. Supervision (to ensure compliance)
4. Reporting (to confirm adherence)

This felt like management. In many organizations, it still feels like management.

It persisted because execution was expensive and correction was slow. If changing direction cost months of rework, then prescribing the path in advance appeared rational. Supervision appeared necessary because the cost of drift was high and the speed of detection was low.

The result was an entire generation of leaders trained to direct rather than orient. To supervise rather than evaluate. To measure compliance rather than outcomes.

AI does not tolerate that pattern.

Give an AI system detailed step-by-step instructions for a complex task and it will follow them — rigidly, literally, without the contextual judgment that a human would silently apply. The output is technically correct and frequently wrong. Anyone who has over-specified a prompt has experienced this.

Give the same system clear intent, well-defined constraints, and freedom to determine approach — and the output improves dramatically. Not because the AI became smarter. Because the human communicated more effectively.

The same principle applies to people.

It always has.

The Mechanism

The structural equivalence is not a coincidence. It emerges from a shared condition: both human teams and AI systems operate under uncertainty with incomplete information.

No leader possesses perfect knowledge of what their team will encounter during execution. No prompt author possesses perfect knowledge of how an AI system will interpret ambiguity. In both cases, the executor must make micro-decisions that the director cannot anticipate.

When the director prescribes every step, two things happen:

First, the executor loses the ability to adapt when reality diverges from the plan. They follow the instructions even when the instructions no longer make sense, because compliance was the expectation.

Second, the director becomes a bottleneck. Every unanticipated situation requires escalation, because the executor was never given the context to exercise judgment — only the steps to follow.

Intent-based direction solves both problems simultaneously. The executor understands the destination and the boundaries. When reality presents an unanticipated situation, they can make a judgment call that advances the intent without violating constraints. No escalation required. No bottleneck created.

This is why the best teams and the best AI interactions share the same structure. Both succeed when the human provides orientation rather than instruction.

What AI Revealed

Organizations have debated management philosophy for decades. Command-and-control versus empowerment. Centralization versus decentralization. Trust versus oversight. These debates never fully resolved because both sides could point to failures of the other approach.

AI collapses the debate.

It does so because AI provides a controlled experiment that organizations never had with people. When you over-specify a prompt, you can watch the AI follow your instructions into absurdity — in seconds rather than months. When you provide clear intent with good constraints, you can watch the AI discover creative solutions you did not anticipate — also in seconds.

The feedback loop that takes months with human teams happens in moments with AI. The lesson becomes undeniable.

Over-specification produces brittle execution.

Intent produces adaptive execution.

This is not new knowledge. It is old knowledge made viscerally obvious.

Every leader who has learned to prompt AI well has — whether they realize it or not — been practicing the leadership model that their organization always needed. The question is whether they recognize the implication: if this is how you get the best work from AI, it is also how you get the best work from people.

The Uncomfortable Implication

If the intent-constraints-autonomy-evidence-correction loop is the model that works for both AI and people, then a significant amount of organizational management infrastructure exists not because it produces better outcomes, but because leaders were never trained to operate any other way.

Status meetings exist because leaders cannot articulate intent clearly enough to trust outcomes.

Detailed task assignments exist because organizations lack the constraints that would make autonomous execution safe.

Approval gates exist because evidence-based correction was never built into the operating rhythm.

Micromanagement is not a personality flaw. It is the predictable result of an organization that never invested in the infrastructure of intent.

AI makes this visible because AI forces the conversation. You cannot micromanage a prompt into producing good work on a complex task. You must learn to communicate differently. Organizations that learn this lesson only for AI — while continuing to micromanage people — will find themselves in an absurd position: treating machines with more trust than humans.

The Inversion

The conventional narrative is: "We need to learn how to manage AI."

The more accurate framing may be: "AI is teaching us how we should have been managing all along."

The management model that AI demands — clear intent, explicit constraints, autonomous execution, evidence-based evaluation — is not a new invention adapted from technology. It is the rediscovery of a leadership principle that organizations systematically abandoned as they grew, bureaucratized, and optimized for predictability over adaptability.

AI did not create a new operating model.

AI revealed the operating model that was always there — buried under decades of process designed for a world where execution was expensive and correction was slow.

Why This Still Requires Change

If the operating model is not new, an obvious question follows: why does anything need to change?

The answer is that knowing the correct model and being able to run it are different problems.

Most organizations have understood for decades that intent-based leadership produces better outcomes than micromanagement. The language exists. The business books exist. The case studies exist. Very few leaders would argue in principle that detailed instruction is superior to clear intent.

And yet most organizations still operate through detailed instruction.

The reason is not ignorance. It is that running the intent model requires infrastructure that most organizations never built — because they never had to.

When execution was performed exclusively by humans, the gaps were compensated silently. An experienced engineer received a vague directive and filled in the missing intent from context, institutional memory, and pattern recognition. A senior leader gave ambiguous constraints and relied on the team's judgment to interpret them reasonably. The model appeared to function even without explicit intent, because humans performed the translation work invisibly.

AI does not compensate.

A machine given vague intent produces vague output. A machine given no constraints operates without boundaries. A machine given no success criteria optimizes for whatever metric is easiest to maximize. Every gap in organizational clarity that humans once papered over becomes a visible failure when machines execute.

This is why organizations must change — not because the model is wrong, but because the model was never fully implemented. What changes is not the philosophy. What changes is the rigor.

Specifically, organizations must now invest in:

Explicit intent. Not "we all know what we're trying to do" but written, testable declarations of desired outcomes that both humans and machines can evaluate against. Intent must become an artifact, not an assumption.

Codified constraints. Not constraints that live in a senior architect's head, but constraints expressed precisely enough that an autonomous system can respect them without asking. Architectural principles, security boundaries, ethical limits, and financial guardrails must exist in a form that scales beyond human memory.

Continuous evidence loops. Not quarterly business reviews, but rapid feedback mechanisms that detect drift early enough to correct it cheaply. When execution happens in hours rather than months, evaluation must happen in days rather than quarters.

Trust as a system property. Not trust as a cultural aspiration — "we trust our teams" — but trust as a measurable outcome of demonstrated capability. Organizations must build the mechanisms that allow autonomy to expand as evidence of sound judgment accumulates.

None of this is philosophically new. All of it is operationally hard.

The change is not adopting a different model. The change is finally doing the difficult work that the correct model always required — work that was previously optional because humans compensated, and is now mandatory because

machines will not.

AI did not change what good leadership looks like.

It raised the cost of pretending to practice it.

Implications

For leaders: The skill of prompting AI well is the same skill as leading people well. If you find yourself writing better prompts than you write team charters, the gap is not in your AI literacy. It is in your leadership practice.

For organizations: The governance infrastructure built for AI (intent documents, constraint definitions, evaluation criteria, feedback loops) is the same governance infrastructure that human teams need. Build it once. Apply it everywhere.

For AI adoption: Organizations struggling with AI adoption may not have a technology problem. They may have a leadership clarity problem that AI is exposing for the first time. If you cannot express clear intent to a machine, you probably were not expressing clear intent to your people either — they were just compensating silently.

Open Questions

If the operating model is the same for people and machines, does the organization need separate governance structures for AI execution and human execution — or should they converge?

What happens to middle management roles that exist primarily to translate strategy into detailed instructions? If intent can flow directly to capable executors (human or AI), what is the new value proposition of the translation layer?

Is there a risk that organizations learn to lead AI well while continuing to lead people poorly — creating a two-tier system where machines receive better direction than humans?

Principle

AI did not invent the operating model that produces the best outcomes. It made that operating model impossible to ignore.

This paper is part of Project Intent, an evolving body of work exploring how intelligent organizations should see, decide, and execute when intelligence is abundant.