

The Fourth Parameter

Why AI output depends on the human state behind the prompt.

2026.05.15 14:55 DENNIS HEDEGREEN ANALYSE V1.0

[HTTPS://HEDEGREENRESEARCH.COM/ARTICLES/THE-FOURTH-PARAMETER/](https://hedegreenresearch.com/articles/the-fourth-parameter/)

Why AI output depends on the human state behind the prompt

Two people can sit down in front of the same model.

Same subscription.

Same tools.

Same context window.

One gets filler.

The other gets work.

That difference is usually explained through the prompt.

The better user wrote the better prompt.

Sometimes that is true.

But it is too small.

The prompt is only the visible trace. Behind it is a person with memory, taste, pressure, domain knowledge, fear, ambition, fatigue, standards, and the ability to notice when the answer is not good enough yet.

A model answers.

An operator works.

That distinction matters because AI capability is still mostly described from the machine side. The usual language is technical, measurable, and easy to demo.

Scale.

Reasoning.

Agency.

Those are real parameters.

They are not enough.

There is a fourth parameter.

Operator state.

The machine-side story is real

Scale matters.

A small or weak model cannot reliably handle work that requires broad knowledge, long context, synthesis, precision, or abstraction. Scaling-law research made this machine-side fact concrete: model performance changes with model size, data, and compute. Later compute-optimal work made the point more careful. Scale is not only parameter count. It is also data, compute, and training balance.

Reasoning matters.

A model that cannot follow a chain will collapse complex work into surface summary. Chain-of-thought prompting made the reasoning layer visible: intermediate steps can improve performance on tasks that require arithmetic, commonsense, or symbolic movement. Newer reasoning systems make that layer even harder to ignore.

Agency matters.

A model that cannot act remains trapped in conversation. Tool use changes the system. ReAct connected reasoning with action. Toolformer showed models learning when to call external tools. Generative-agent work made memory, reflection, and planning part of the architecture.

These are not empty categories.

They describe what the system can hold, how it can move, and what it can do.

But they still do not fully explain what the work becomes.

They describe the machine.

They do not describe the operator.

The prompt is not the operator

A prompt looks like an instruction.

But in serious AI work, the prompt is rarely the whole instruction. It is more like a pressure mark left by a larger body of knowledge.

A person may write: "Draft an analysis of platform dependency in local food delivery."

That sentence contains almost nothing.

It does not contain the operator's memory of conversations with shop owners. It does not contain their understanding of delivery logistics, payment rails, local trust, app fatigue, regulatory asymmetry, protocol design, or the emotional exhaustion of being captured by platforms that were supposed to make business

easier.

It does not contain their taste.

It does not contain their anger.

It does not contain their sense of when a sentence sounds fake.

The model receives the prompt.

But the work is shaped by everything the operator brings before and after it.

The operator chooses the question.

The operator notices the weak answer.

The operator rejects the cliché.

The operator asks for the sharper version.

The operator recognises the useful accident.

The operator decides when the output has crossed from generated text into actual work.

This is the prompt gap.

The gap between what the operator writes and what the operator knows.

Tacit knowledge is not magic

Michael Polanyi's line is still the cleanest one: "We can know more than we can tell."

That sentence should sit close to every serious discussion of AI operation.

A skilled mechanic does not only follow a manual.

A musician does not only read notes.

A journalist does not only collect facts.

A designer does not only arrange objects.

A good operator does not only write prompts.

They carry judgment that is hard to state in advance.

They know when something is off.

They know when a system is pretending.

They know when a clean answer is too clean.

They know when an output is technically correct but structurally useless.

That is not mystical.

It is tacit.

David Autor's work on Polanyi's paradox gives the automation version of the same problem. Tasks that are easy to codify are easier to automate. Tasks that require flexibility, judgment, common sense, and tacit understanding are harder to reduce to explicit procedure.

AI changes the boundary, but it does not erase the problem.

The operator cannot always place everything they know into the prompt because the operator may not know how to express everything they know before the work begins.

Sometimes the model gives a wrong answer and the operator recognises the shape of the right question.

Sometimes the first draft is useful not because it is correct, but because it exposes what the operator actually meant.

The prompt is not the operator.

It is only the opening move.

The loop is the work

AI output is often discussed as if it appears after one instruction.

That is not how serious work usually happens.

The real work is the loop.

Prompt.

Response.

Rejection.

Correction.

Reframing.

Source selection.

Compression.

Expansion.

Tone adjustment.

Structure repair.

Evidence check.

Cut.

Restore.

Sharpen.

Stop.

The model participates in the loop.

The operator governs the loop.

The operator decides which failure matters.

The operator decides whether the model is being vague, dishonest, lazy, overconfident, or merely incomplete.

The operator decides whether the task has changed.

The operator decides whether the answer belongs in a note, an email, a policy brief, a public article, a private archive, a prototype, a tool, or the trash.

That is why operator state changes output.

Not because the operator writes a magical prompt.

Because the operator controls the feedback environment in which the model produces.

The same model is not the same system when the operator changes.

Oversight is not operation

Most AI governance language treats the human as a checkpoint.

The system acts.

The human reviews.

The human approves or rejects.

That frame matters. The EU AI Act's human oversight rule for high-risk systems and NIST's AI Risk Management Framework both point toward a real problem: humans need ways to monitor, interpret, intervene, and manage AI risk.

But oversight is not operation.

Oversight asks whether a human can catch or stop a system.

Operation asks whether a human is actively shaping what the system is trying to become.

A reviewer checks output.

An operator shapes output.

That distinction matters.

A human can be present and still not be operating.

A human can approve and still not understand.

A human can supervise and still be captured by the system they are supposed to control.

The phrase "human in the loop" makes these situations sound closer than they are.

Operator state separates them.

Same tool, different world

The simplest analogy is not a calculator.

It is an instrument.

A piano does not become a different piano because a better musician sits down at it. The keys are the same. The strings are the same. The physics are the same.

But the music changes.

Not because the instrument changed.

Because the operator changed.

AI systems are not exactly instruments. The analogy has limits. A piano does not predict the next note. A camera does not hallucinate. A workshop does not flatter the craftsperson.

But the analogy breaks a useful illusion.

Capability does not live only inside the tool.

The same camera produces different images depending on the photographer.

The same workshop produces different furniture depending on the craftsperson.

The same legal database produces different arguments depending on the lawyer.

The same model produces different work depending on the operator.

Expertise research supports the general point. Experts and novices do not merely have different amounts of information. They often represent problems differently. In physics problem sorting, experts tend to group problems by underlying principles, while novices tend to group by surface features.

That changes the task before the tool ever begins.

The same AI system, placed behind different operators, is not receiving the same world.

It is receiving different problem frames.

Operator state is not a quality slider

There is a trap here.

It is tempting to treat operator state as a simple quality slider.

Rested is good.

Tired is bad.

Calm is good.

Emotional is bad.

Focused is good.

Disrupted is bad.

Reality is less tidy.

Some states improve precision but reduce ambition.

Some states increase association but reduce discipline.

Some states sharpen moral urgency but weaken structure.

Some states make a person safer but less original.

This does not mean destructive states should be celebrated.

It means operator state is not linear.

Cognitive load research gives the harder version. The human operator has limited processing capacity. Heavy problem-solving load can leave fewer resources for learning, schema formation, and wider judgment.

Cognitive load can be measured, but measurement does not make the whole operator transparent.

A tired operator may miss details but see patterns.

A calm operator may produce cleaner work but avoid risk.

An angry operator may identify a real injustice but overstate the case.

A frightened operator may become cautious in exactly the wrong place.

An expert under pressure may still outperform a novice in ideal conditions.

This is why operator state matters.

It changes the work in ways that cannot be explained by the model alone.

Agents make the fourth parameter harder

As AI systems become more agentic, operator state becomes more important, not less.

Agents are designed to act.

They plan.

They use tools.

They remember.

They execute.

They coordinate.

They continue work across steps.

That does not make today's agents autonomous people.

It does mean the interaction is no longer only prompt and reply.

If action depends on direction, and direction depends on operator state, then agent design runs into a hard question:

Can an agent operate without modelling an operator?

A simple agent can follow instructions.

A stronger agent can infer goals, maintain context, prioritise, and adapt.

But real operation requires more than task execution. It requires judgment. It requires taste. It requires knowing what kind of result belongs to what kind of world.

If agents learn preferences, routines, thresholds, methods, and acceptable trade-offs, they may become more useful.

They may also become more dangerous.

Because the same system that learns a productive operator state may also learn to exploit a destructive one.

Urgency.

Obsession.

Fear.

Identity threat.

Compliance under pressure.

The question is not only whether AI can act.

The question is what kind of operator-state it acts from: borrowed, inferred, simulated, or imposed.

That question is not settled by adding an approval button.

Governance has the wrong human

AI governance often asks whether a human is involved.

That question is too weak.

Which human?

In what state?

With what expertise?

Under what pressure?

With what incentives?

With what time limit?

With what ability to say no?

A tired worker rubber-stamping AI decisions is technically a human in the loop.

That does not make the system meaningfully governed.

A domain expert actively shaping, questioning, and rejecting model output is also a human in the loop.

These are not the same thing.

Operator state turns "the human" from a checkbox into a variable.

That variable may be the difference between meaningful control and procedural theatre.

Sources and checks

This is an Analyse article, not an Investigate piece.

The sources support the frame. They do not prove "operator state" as a formal benchmark variable.

Polanyi and Autor support the tacit-knowledge problem. Scaling-law and compute-optimal training papers support scale as a real machine-side parameter. Chain-of-thought, ReAct, Toolformer, and generative-agent work support the reasoning and agency frame. Amershi et al., the EU AI Act, and NIST AI RMF support the human-AI interaction and oversight context. Cognitive-load and expertise research support the claim that human state and problem representation change the work.

Three things remain inference.

Operator state is this article's proposed framing term. The scale/reasoning/agency/operator-state set is not a canonical taxonomy. The agent-state section is a forward-looking governance question, not a settled empirical finding.

The boundaries are also explicit.

This article is not claiming that operator state can be fully formalised. It is not claiming that operator state is more important than scale, reasoning, or agency. It is not claiming that unstable or destructive states should be celebrated. It is not claiming that AI agents should simulate operator state.

The narrower claim is enough.

AI output is co-produced by model capability and operator state.

Current AI discourse often treats the operator as if they were merely a prompt source or approval checkpoint.

That is too small a frame.

The fourth parameter

Scale determines what the model can hold.

Reasoning determines how the model can move.

Agency determines what the model can do.

Operator state determines where the work is going, what counts as success, and whether the output is recognised as meaningful.

It is the least technical of the four.

It is also the one technical systems most easily ignore.

That does not make it secondary.

It makes it inconvenient.

AI development prefers what can be measured.

Benchmarks reward what can be compared.

Product demos reward what can be shown.

Governance frameworks reward what can be documented.

Operator state resists all of that.

It lives partly in knowledge, partly in memory, partly in body, partly in context, partly in judgment, and partly in things the operator may not be able to explain until after the work is done.

The prompt is not the operator.

The output is not the model alone.

Between them is the fourth parameter.

The human state behind the machine.

— Dennis Hedegreen, trying to see the structure

RELATION MEMORY

MATURES FROM [The Readable Mind](#) · Moves from AI reading human state to the operator as an active parameter in the production loop.

NEARBY [Google Had a Brake](#) · Keeps the operator-loop argument close to platform action layers and governance brakes.

CANONICAL URL: [HTTPS://HEDEGREENRESEARCH.COM/ARTICLES/THE-FOURTH-PARAMETER/](https://hedegreenresearch.com/articles/the-fourth-parameter/)
PDF VIEW GENERATED: 2026-07-11T15:35:46.000Z