

The Judgment Gym

Dømmekraft — and whether your AI is a simulator or a stimulator

Start with the word

The Danish **dømmekraft** says something the flat English “judgment” loses. It is Kant's *Urteilkraft* — the **power** of judgment — and it names a specific faculty: the ability to bring a particular case under the right general idea **when no rule tells you how**. Aristotle called the practical version *phronesis*. Its defining property is that it is **non-algorithmic by nature**. If a procedure could produce the answer, you would not need judgment — you would need a calculation.

Dømmekraft is exactly the faculty you reach for when the rule runs out: weighing a spread of clashing, non-comparable evidence — a statistic here, an anecdote there, a base rate, a gut-level red flag, someone's passionate testimony — none of which add up mechanically, into one coherent read. That synthesis across things that do not share a common unit is the human act this module is about. And it is **tacit**: as the philosopher Michael Polanyi put it, *we know more than we can tell*. You cannot hand someone your dømmekraft as a list of steps.

Start with an experiment on yourself

A friend asks whether they should take a new medication. Here is everything you know. Form a judgement — then watch how you did it:

- **Statistic**: in trials, it helped 6 in 10 people with this condition.
- **Base rate**: the condition often improves on its own within months.
- **Anecdote**: your friend's cousin had a bad reaction to it.
- **Red flag**: the prescribing clinic also sells the drug.
- **Testimony**: a doctor your friend trusts recommends it.

You arrived at some lean — cautious, or supportive, or “get a second opinion.” Now try to write the **formula** you used. How many anecdote-points equal one base-rate-point? What exchange rate converts a red flag into a percentage? **There isn't one**. You did not compute; you *judged* — you held incommensurable things together and felt your way to a read you cannot fully justify in steps. That felt, unformalisable synthesis is dømmekraft. Hold onto how it felt, because this module is about what happens to that faculty when a machine offers to do it for you.

The two words: simulation and stimulation

A single letter separates two completely opposite ways to use AI on a judgement.

Simulation is when the AI performs the judgement *for* you. It hands you the finished verdict — the decision, the synthesis, the “on balance, X.” The faculty fires inside the machine (or appears to), and you receive the output, clean and confident.

Stimulation is when the AI provokes *your* judgement — it feeds you the counter-argument, surfaces the datapoint you skipped, names the assumption you did not know you were making, gives you the base rate you ignored — but leaves the **act of judging to you**. Same tool, two modes. One

substitutes the faculty; the other exercises it.

And here is the sharp end, the part that matters most for you specifically right now: **judgement is a muscle, and simulation lets you skip the workout while feeling like you did it.** You got the answer, the essay, the tidy synthesis — so it *feels* like learning happened. But the faculty never contracted. No reps. This is your whole course objective turned inside out: you learn by effortful sense-making, and simulation is an **effort-removal machine that disguises itself as a productivity gain.**

Worse, generative AI **defaults** to simulation. It is built to emit finished, fluent output — that is its factory setting. Which means its real educational value, stimulation, runs *against its own grain*. You have to deliberately steer it there. Left to itself, it will always offer to do your thinking for you, pleasantly.

The durable insight (write this down)

Ask whether the AI is simulating your judgement or stimulating it — doing the reps for you, or making you do them. Simulation hands you a verdict and quietly retires the faculty. Stimulation withholds the verdict and exercises it. GenAI defaults to simulation; learning lives in stimulation. The relief you feel when it just answers is the exact feeling of a workout skipped.

The honest part: sometimes you should let it decide

If this module told you to always trust your own dømmekraft, it would be lying to you — and the earlier modules were about not being lied to, including by your teachers. There is a large body of research, from Paul Meehl in 1954 through Daniel Kahneman's recent work on “noise,” showing that in **repeated, well-structured prediction tasks**, a simple validated model routinely **beats** unaided expert intuition. Human judgement is noisy, overconfident, and biased in exactly the places we most trust our gut. So “follow your instinct” is not just naïve here; it is often empirically wrong.

You may recognise this boundary from the world of measurement: the entire premise of a structured assessment is that disciplined measurement can outperform an interviewer's hunch. Every serious instrument faces this exact design choice — whether to *simulate* the judgement (“this candidate is a poor fit”) or to *stimulate* it (“here is structured evidence for a human to weigh”). Neither answer is automatically right. One caution worth carrying from the bias module: a validated model is still a lens ground on some particular data. Deferring to it is often wise — *and* it still deserves the question “whose data built this, and where might it not apply?”

The one judgement you can never outsource

The module's real thesis is not “simulation bad.” It is this: **the judgement of WHEN to defer to the simulation and when to override it is itself a judgement — and that one you can never hand over.** Because to outsource it, you would have to ask the simulation whether to trust the simulation. That meta-dømmekraft — matching the mode to the situation — is the irreducible human core the defensive half of this course has been circling: the faculty that does all the seeing-through.

A rough guide to which mode fits

Lean toward letting it decide (simulate) when the task is repeated, well-structured, and has a validated model behind it, and when the stakes of a single call are low and checkable. **Lean toward**

doing your own judging (stimulate) when the decision is novel, one-off, values-laden, or when the whole point is to *learn* — because there, the reps are the reward, and a borrowed verdict is the booby prize.

The pivot: trace versus act

This module is the turning point of the course, because it names what was hiding under the first three. An AI produces the **linguistic trace** of a judgement — “there are compelling arguments on both sides; weighing the evidence, the most defensible conclusion is...” — the full grammar of wise deliberation, **without the act of deliberating**. That is the same gap you have met three times already:

MODULE	THE SURFACE / TRACE	THE ACT IT LACKS
1 · Reading vs. Meaning	fluent text (tokens)	actually meaning it
2 · The Reference Chain	perfect citation (form)	actually pointing at something
3 · The Lens & the Smudge	objective register	actually having a standpoint it owns
4 · The Judgment Gym	the language of deliberation	actually deliberating — dømmekraft

The first three modules taught you to distrust the surface. This one names what the distrusting is *done with* — dømmekraft — and warns that the same tool can quietly retire it while you applaud the output. That closes the **defensive** arc of the course: seeing through fluent text, fake citations, hidden standpoints, and borrowed judgement. What comes next turns outward — from seeing through answers to **aiming at the right questions**: finding the real problem, and building understanding across a whole conversation rather than one prompt. The faculty you've just met is what does the aiming.

How to prompt for stimulation, not simulation

INSTEAD OF “What should I do?” / “Just give me the answer.”
ASK “Argue the strongest case against my conclusion.”
ASK “What evidence am I underweighting or missing?”
ASK “What's the base rate here, and am I ignoring it?”
ASK “Name the assumptions my reasoning depends on.”
ASK “Give me the questions I'm not asking – but keep the verdict.”

The four-question drill

- 1 · WHO'S JUDGING? Is the faculty firing in me – or only in the machine?
- 2 · SIMULATE / STIMULATE? Did I ask for the verdict, or for a better fight?
- 3 · DEFER OR OWN IT? Repeated, structured, validated model → defer. Novel or values-laden → own it.
- 4 · DID I DO THE REPS? If I feel relief, did I just skip the workout?

Try it yourself: the companion widget

The **Judgment Gym** deals you real decisions with messy, clashing evidence. You commit to a call, then get to feel both modes — ask the AI to *decide*, and ask it to *challenge* you — and finally judge which mode the situation actually deserved. Some reward doing your own thinking; at least one rewards deferring to the model. Learning to tell them apart is the whole skill.

Reflection questions

1. Recall the medication experiment. If an AI had simply told you “cautiously supportive,” would you have noticed everything you weighed to get there? What would you have lost by taking the verdict?
2. Where in your studies right now are you using AI as a simulator when you'd learn more using it as a stimulator? Be specific — which task, and what would the stimulating prompt be?
3. Describe a decision where deferring to a validated model is wiser than trusting your gut. How did you know? What would have to change to make it a “own it” decision instead?
4. You can't outsource the judgement of when to trust the AI. What is your personal rule of thumb for that meta-decision — and how will you keep testing it?

A note for right now: you are in a rare window — the years when your main job is to build the faculty itself, not just spend it. Every time you let AI simulate a judgement you could have made, you spend that window; every time you make it stimulate you instead, you invest it. The tool is not the enemy. Using it against its own grain — to make yourself think harder, not less — is what an AI-confident mind does. Test that claim too.