

The Scope Finder

How to widen a problem without blurring it — and where AI actually helps

The tension the whole module lives on

Two failures sit at opposite ends of a rope, and problem-finding is the act of not falling off either end. **Tunnel vision:** you grab the first tidy version of the problem and solve it fast — and it was the wrong problem, tidy precisely because it was too narrow to be true. **Fog:** you open the problem up, everything connects to everything, a hundred framings bloom, and you can no longer act — analysis paralysis dressed up as open-mindedness.

Here is the trap in one line: **curiosity without discipline is noise; discipline without curiosity is blindness.** The skill is not choosing a side. It is a **rhythm** — breathe out, then breathe in. **Expand** the scope wide enough to be sure you are aimed at the real thing; then **contract** it down to something you can actually grip. And the punchline for this course: AI is unusually strong on the expand stroke and quietly useless — sometimes harmful — on the contract stroke. Knowing which stroke you are in is the literacy.

Start with an experiment on yourself

A student says: “*My problem is that I need to stop procrastinating on my thesis.*” Before reading on, notice: that already feels like a problem, and the solution feels obvious (more discipline, an app, a schedule). Now widen it once — ask “**what would that get me?**”

- I want to stop procrastinating ... **so that?**
- ... I make progress on the thesis ... **so that?**
- ... I stop feeling anxious every evening ... **so that?**
- ... I can trust myself to finish things. ← *there it is.*

In four steps the problem moved from “install a scheduling app” to “I don’t trust myself to finish” — a completely different problem with completely different solutions. Notice two things. First, the original statement had a **solution smuggled inside it** (“stop procrastinating” presumes the fix is willpower). Second, you could keep laddering forever — so the skill is not only expanding, it is knowing **when to stop and commit**. That is the contract stroke, and it is the harder half.

What a problem actually is

A workable definition, because each part is a **test** you can run: a problem is a **gap** between how things are and how they should be, **that matters**, and that **isn’t yet resolved**. If you can’t name both the current state and the desired state, you don’t have a problem yet — you have a feeling. If closing the gap wouldn’t change any outcome you care about, you have a *defect*, not a problem. And a gap only exists against a “should” — which always belongs to **someone**. Change whose “should” you take, and the problem changes shape.

The expand stroke: curiosity as repeatable moves

Curiosity is not a personality you were or weren't born with. It is a set of moves you can run on any situation — and AI, having read more contexts than any person, lowers the effort of every one. Use it as a **breadth engine**:

Multiply the stakeholders. “Who else touches this, and what does the problem look like from each of their chairs?” This externalises the standpoints you are structurally blind to — your own lens, turned into a tool.

Ask for questions, not answers. The most important swap in this whole module: instead of “what's the answer?”, ask “**what are the ten questions I should be asking about this that I haven't?**” That one change turns the AI from an oracle into a curiosity prosthesis.

Ladder the scope. “State this one level wider; now one level narrower.” Wider finds the real goal; narrower finds the part you can move this week.

Borrow from another field. “How is a problem like this solved in a completely different domain?” The fastest way to break a stuck frame.

The durable insight (write this down)

AI expands the space of questions; you contract it to the one worth answering. Divergence is assisted — let the machine open the aperture. Convergence is human — keep your own hand on the focus ring. Scope without blur is exactly this division of labour: the AI widens, you choose.

The contract stroke: where AI can't save you

Expansion feels productive — options fanning out, everything suddenly connected. But an expanded problem you never narrow is just fog with good intentions. Three human moves pull it back to clarity, and none of them can be outsourced:

1 • The relevance filter. Of all these framings, which actually connect to an outcome you care about? Most won't. **Dropping them is the skill** — not a failure to appreciate them.

2 • The “so what” test. For each new viewpoint, ask: if this were true, would it change what I do? If not, it is *interesting*, not *relevant* — and interesting-but-irrelevant is precisely what blurs perception. This is the single most useful filter in the module.

3 • Commit to a frame. Expansion has to end in a choice. You can hold a frame provisionally and still commit to it; refusing to choose is not open-mindedness, it is fog. And this call — which expanded viewpoint actually *matters* — is judgement, the faculty from the last module. The AI can widen the options; it cannot make this cut for you.

The move that separates scope from blur

After every expansion, run the “**so what**” test and **throw most of it away**. A good problem-finder generates ten framings and discards eight without regret. The discard is not waste — it *is* the thinking. If expansion never ends in a ruthless cut, you haven't widened your scope; you've only blurred it.

Defect or problem? A quick test

Not every gap deserves the full treatment, and over-expanding a trivial defect into an existential crisis is its own failure. A **defect** is a gap that fails a spec inside a system that is basically working — fix it and move on. A **problem** is a gap that threatens the *purpose* of the system, or reveals that the spec itself is wrong. The test: does closing this gap change the outcome that matters — or just a metric that was convenient to measure? “The slide’s font is off” is a defect. “No one remembers anything from our slides” is a problem. Spend your curiosity where the second kind lives.

The rhythm, on one line

STROKE	WHAT YOU DO	WHO LEADS
Expand (breathe out)	Multiply stakeholders, generate questions, ladder wider/narrower, borrow frames.	AI assists — it’s a breadth engine
Contract (breathe in)	Relevance filter, “so what” test, discard most of it, commit to one frame.	You lead — it’s a judgement

Prompt patterns for the expand stroke

“Give me five framings of this problem, from five different stakeholders.”
“What are ten questions I should be asking about this that I haven’t?”
“State this problem one level more abstract, then one level more concrete.”
“Is there a solution smuggled into how I stated this? What does it foreclose?”
“How is a problem like this handled in a completely different field?”
“If we solved this and nothing changed, what would the real problem have been?”

Notice every one asks the AI to *widen*, never to *decide*. The moment you ask “so which one should I pick?”, you’ve handed it the contract stroke — the one it can’t do and you must.

Where this sits in the series

The first four modules were **defensive** — seeing through fluent text, fake citations, hidden standpoints, borrowed judgement. This one is the first **generative** module: it is about what happens *before* all of that — making sure you are aimed at a real problem at all. A perfectly solved wrong problem is the most expensive mistake there is, because it feels exactly like progress. And the AI, brilliant at solving, will never once tell you the problem was wrong. That check is yours. The final piece of the course then ties the whole thing together — showing how expanding, contracting, and every skill before it become the moves of a single unfolding conversation, rather than one prompt and one answer.

Try it yourself: the companion widget

The **Scope Finder** gives you student-life situations with a narrow presenting problem. You feel the tunnel, run an expand stroke (an AI surfaces stakeholders, questions, wider and narrower framings) — and then do the hard part by hand: tag every generated viewpoint **keep** or **drop** with the “so what” test, and commit to one sharp problem. One case will tempt you to over-expand a simple defect; resisting that is part of the skill.

Reflection questions

1. Take a “problem” you’re carrying right now and ladder it with “what would that get me?” three times. Where did you land — and was there a solution smuggled into your first version?
2. Which is your default failure — tunnel vision or fog? What does that tell you about whether you should reach for the AI to expand, or force yourself to contract?
3. The “so what” test discards interesting-but-irrelevant framings. Recall a time being interested led you away from what mattered. How would the test have caught it?
4. Divergence is assisted; convergence is human. Write one sentence you could say to yourself to know when to stop asking the AI for more and start choosing.

A note on honesty: an AI cannot tell you which problem matters — it has no stake in your outcome and no access to what you actually care about. It can widen your view generously and surface questions you’d never reach alone, and that is a real gift. But every framing it offers is a candidate, not a verdict. The aperture is its job; the focus is yours. Test that division of labour on your own problems this week.