

Using AI in Research Workflows

Germain Gauthier

Bocconi University

September 2026

How do I use AI?

- AI is part of pretty much everything I do on my laptop.
- Some use cases:
 - Getting a quick overview of papers I am unfamiliar with
 - Writing code, organizing project folders, updating results in a paper
 - Proof-checking that the code, the paper, and the slides are consistent
 - Discussing ideas and analyses when I am working on a project
 - Improving the clarity of my writing and presentations
- I use pen and paper when I need to fully understand something (e.g., theory, statistics) or for mental clarity (e.g., making a to-do list, planning ahead).

How do I not use AI?

- The one thing I do not do with AI: “Write me a full paper and a full slide set on this vague idea while I go play tennis.”
- First, current AIs are just terrible at this when left unattended, so this is a very iterative process, and you need to be in the loop.
- Second, you will be ultimately presenting the paper and the results, so being in the loop allows you to *understand* what you are doing.
 - The journey is almost as important as the final destination in research.
- As mentioned above, AI can still help speed things up (and by a lot).

My setup

I connect the agent to 5 tools:

Dropbox	Project files and data
GitHub	Code history and repository
Overleaf	Paper and slides
HPC cluster	Heavy computation
Gmail	Co-author comments

The project folder

Dropbox/my-project/

-- overleaf/	Local paper copy, tables, and figures
-- code/	Code repository
-- data/	Raw and processed data
-- notes/	Working notes

Claude Code / Codex runs from the project folder; For Claude, I talk to it in the browser (https://claude.ai/code/*) so I can easily visualize PDFs and PNGs.

I always keep the local paper and the online Overleaf project in sync.

Principle 1: You are in charge.

- Agents can overproduce, expand the scope, and overreact to details.
- I don't let agents run unattended for long periods of time (e.g., overnight) unless the task is very well-defined (e.g., it has to monitor some jobs on the HPC and relaunch them if they fail).
- I have a hands-on approach to my sessions with AI. I stay in the loop and often have to correct course or make judgment calls.
- Being hands-on also ensures I still understand what I am doing, and that I can map it properly in my head though I am not doing the coding anymore.

Principle 2: Prioritize accuracy, not speed or cost-cutting.

- I run models at maximum reasoning. For EVERYTHING.
- A mistake in a paper can cost far more than extra tokens or waiting time, and the difference in IQ between reasoning levels is very large.
- More reasoning does not remove the need for verification (more on this after).

A typical loop

1. **I specify the task in Overleaf.**

Sample, equation, estimator, clustering, and expected output. Just like I would in a detailed PAP or when writing a paper.

2. **A first agent implements it.**

I ask it to flag ambiguities before coding. There is very high value in making sure the AI and you are on the same page.

3. **Other agents review it in a fresh context.**

It reads the specification, code, and outputs; identifies discrepancies between the Overleaf and the implementation.

A typical failure mode is AI making consequential decisions without asking (e.g., dropping some observations). I always ask the agents to check for this, too.

4. **I decide what to revise or accept.**

I resolve issues, rerun, and repeat the review until I'm satisfied.

Two phases: prototyping, then sharpening

Prototyping phase:

- It is an exploratory phase.
- I do not review the code or the prose.
- I treat these outputs as probable but provisional.

Sharpening phase:

- If I think there is something there, I get involved in the writing and coding process.
- I audit the code, paying special attention to the number of rows in the dataset and results that “look weird”. I have many independent agents check the code.
- I rewrite the Overleaf and work with AI to iterate on the visuals.
- I take responsibility for the results and claims.

Communication with co-authors

I keep three distinctions very explicit:

Origin	What the AI produced; what I produced.
Status	What is exploratory; what is polished.
Responsibility	What I have not checked; what I am ready to stand behind.

Your co-authors will not look at “exploratory” and “polished” results in the same way, and won’t make the same comments.

AI use and understanding vary across co-authors. Though I think everyone will eventually use AI and that it greatly reduces coding mistakes and inconsistencies, it’s important to be transparent and to take responsibility for your work.

Using AI in Research Workflows

Germain Gauthier

Bocconi University

September 2026