



WHAT ARE WE MISSING? · A NOTE FROM THE AUTHOR

Cognitive offloading.

Every day we hand a little more of our thinking to machines. Search remembers for us. Maps navigate for us. And now a machine will draft your analysis, summarise your evidence, check your reasoning and offer its confidence for free. Each handoff is convenient, and each one makes the next one easier. Nobody decides to stop thinking. It happens one convenience at a time.

The name for this is cognitive offloading: letting something outside your head do the work your head used to do. Used well, it frees the mind for harder problems. Used carelessly, it trains a reflex the research record has now measured.

This is not a sermon against technology; it is on the record. When a machine's answer arrives wearing authority, even experts follow it into error - and the studies watched it happen. In a 2023 experiment, 27 radiologists read mammograms with an AI assistant that sometimes suggested the wrong category: when the suggestion was wrong, accuracy on those reads collapsed to 20 per cent among inexperienced readers and 46 per cent even among the very experienced; when it was right, they were around 80 per cent correct (Dratsch and colleagues, *Radiology*, 2023). In a retrospective reader study, ten NHS breast-screening readers re-read 60 screening examinations with and without AI prompts: on cases where the AI had missed the cancer, median sensitivity fell from 71 per cent unassisted to 39 per cent with the prompts on screen (Taib and colleagues, *Radiology*, 2026). Nor is the hazard specific to scans: in a randomised trial, 44 physicians - all trained in AI literacy, all free to ignore the advice - saw their composite diagnostic-reasoning scores fall by an adjusted 14.0 percentage points when the chatbot's recommendations were deliberately flawed (Qazi and colleagues, *NEJM AI*, 2026).

A slower hazard sits beneath this one. At four endoscopy centres in Poland, adenoma detection in colonoscopies performed without AI assistance fell from 28.4 per cent before routine AI polyp detection arrived to 22.4 per cent three months into using it (Budzyń and colleagues, *The Lancet Gastroenterology & Hepatology*, 2025). One observational cohort cannot prove the machine caused it; it is a warning consistent with deskilling, not evidence that AI caused it. The judge must keep judging, and keep practising the craft being judged.

This book was built with the machine already in the room. Most books on thinking were written before that was true; this one began with it. It does not ask you to refuse the machine. It trains the discipline that keeps you in charge of it: the habit of judging for yourself, exactly when a fluent answer invites you to skip it. The answer can be outsourced. The asking cannot.

I cannot tell you where your own line sits; that is yours to draw. But the line is worth drawing deliberately, before the convenience draws it for you. What goes missing is never the answer. It is the asking.

The rest is yours.

Sources: Dratsch, T., et al. (2023). *Radiology*, 307(4), e222176. Taib, A.G., et al. (2026). *Radiology*, 320, e252590. Qazi, I.A., et al. (2026). *NEJM AI*, doi:10.1056/Aloa2501001. Budzyń, K., et al. (2025). *The Lancet Gastroenterology & Hepatology*, 10(10), 896-903. All four as cited in Chapter 14 of *What Are We Missing?*; the full chapter notes and reference lists are free at what-are-we-missing.pages.dev (Sources and notes).

Zafar Alam

Civil engineer and Chartered Construction Manager (CIOB).

I began writing *What Are We Missing?* in 2021.

zafaralam.author@gmail.com

what-are-we-missing.pages.dev