



The AI helps you learn. It **won't** do your exam.

An online judge where the teacher sets how much AI a session allows — and where allowed AI use is graded, not ignored.



WHILE LEARNING

The assistant explains. It refuses to write the solution.



IN THE EXAM

The teacher picks the level. Where AI is allowed, prompts are graded too.



FOR THE TEACHER

The same models write exercises, verify them, build papers, mark them.

INTRODUCTION

Who actually wrote this code?

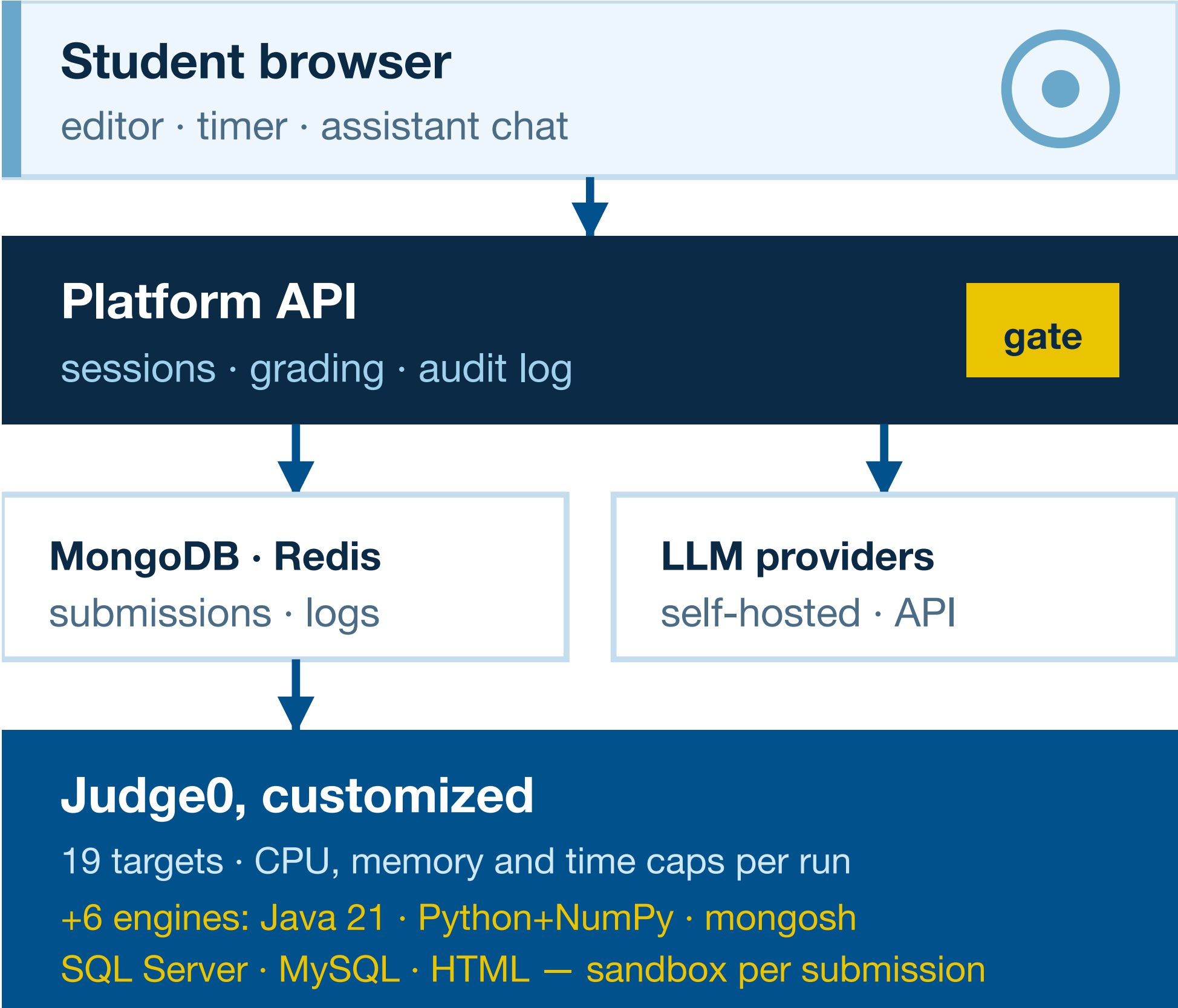
An AI finishes most coursework in seconds. Ban it or ignore it — either way, the grade stops showing who learned.

RESEARCH QUESTIONS

- RQ1** Can "how much AI help" be a configured, auditable setting?
- RQ2** Can an LLM pipeline fill an exercise bank without losing quality?
- RQ3** What must a judge do beyond stdin and stdout?

METHOD · PLATFORM

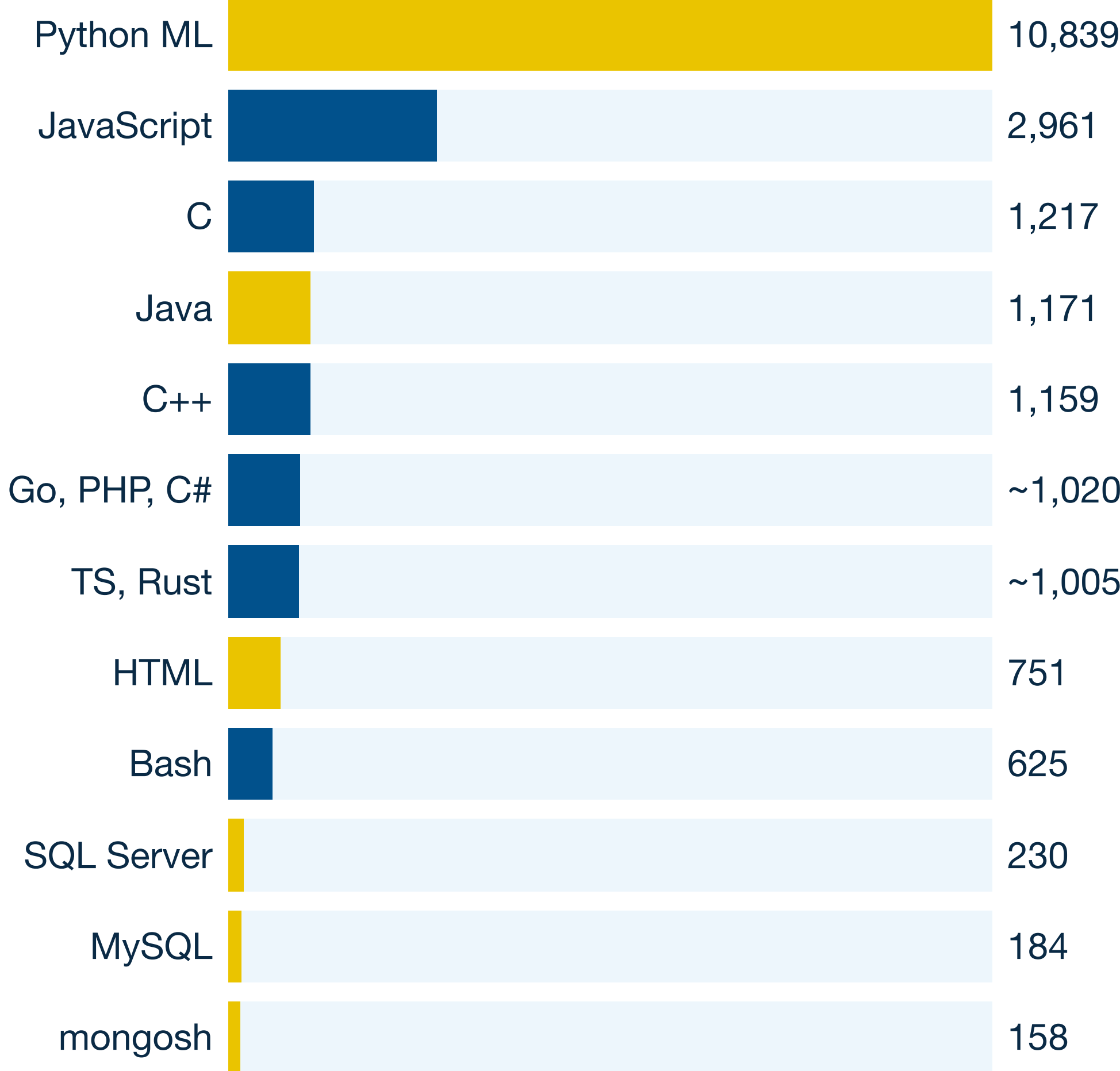
One judge for the whole curriculum



A database exercise gets a real server, seeded and destroyed per submission.

RESULTS · COVERAGE

Where the bank sits



Yellow = the six engines we built on top of Judge0 · blue = languages Judge0 ships with.

METHOD · CONTRIBUTION 1

Help while learning, never the answer

```
1 def total(xs):
2     s = 0
3     for x in xs: s = x
4     return s
```

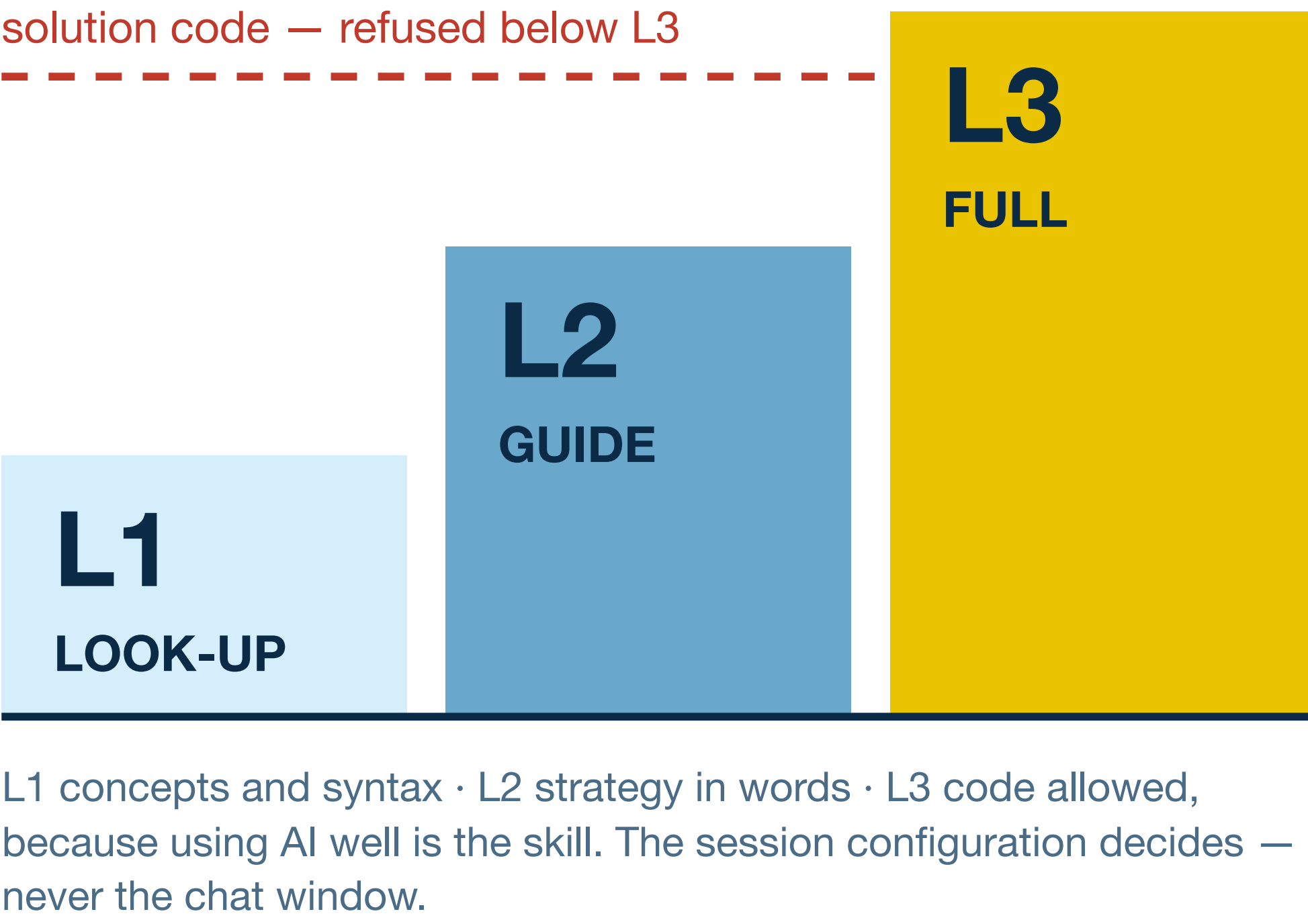
case 1 ✓ ✓ ✗ ✗ score 2 / 4

STUDENT
Just write the correct function for me.

ASSISTANT · LEVEL 2
I won't write the solution. After each step, what should **s** hold — this item, or everything so far?

METHOD · THE LEVELS

The teacher sets the level



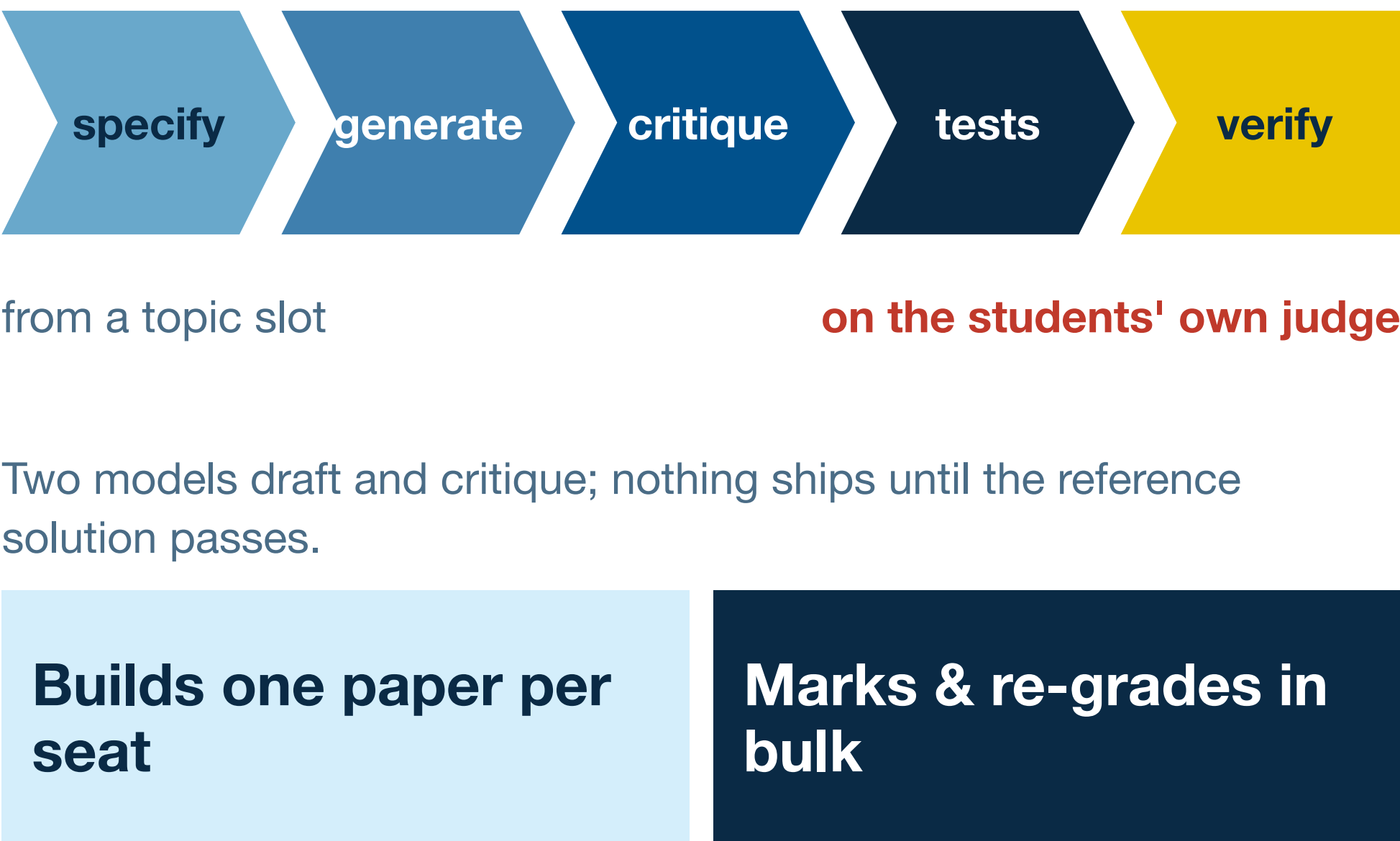
METHOD · EXAMS WITH AI

At L3, the prompt is part of the answer

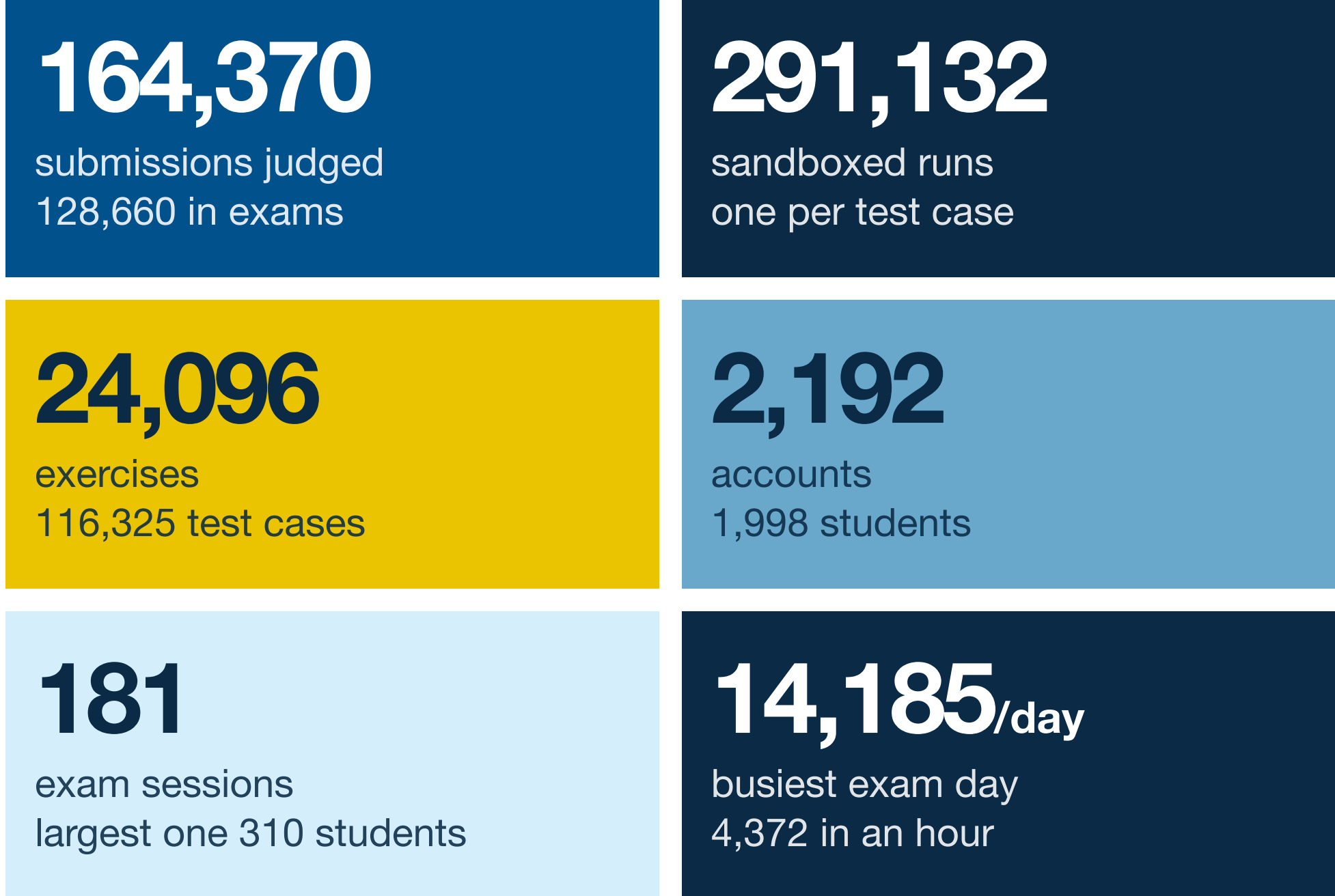
Every turn is stored with the submission. The teacher marks the code the judge scored **and** how the student steered the assistant.

METHOD · CONTRIBUTION 2

The same AI works for the teacher

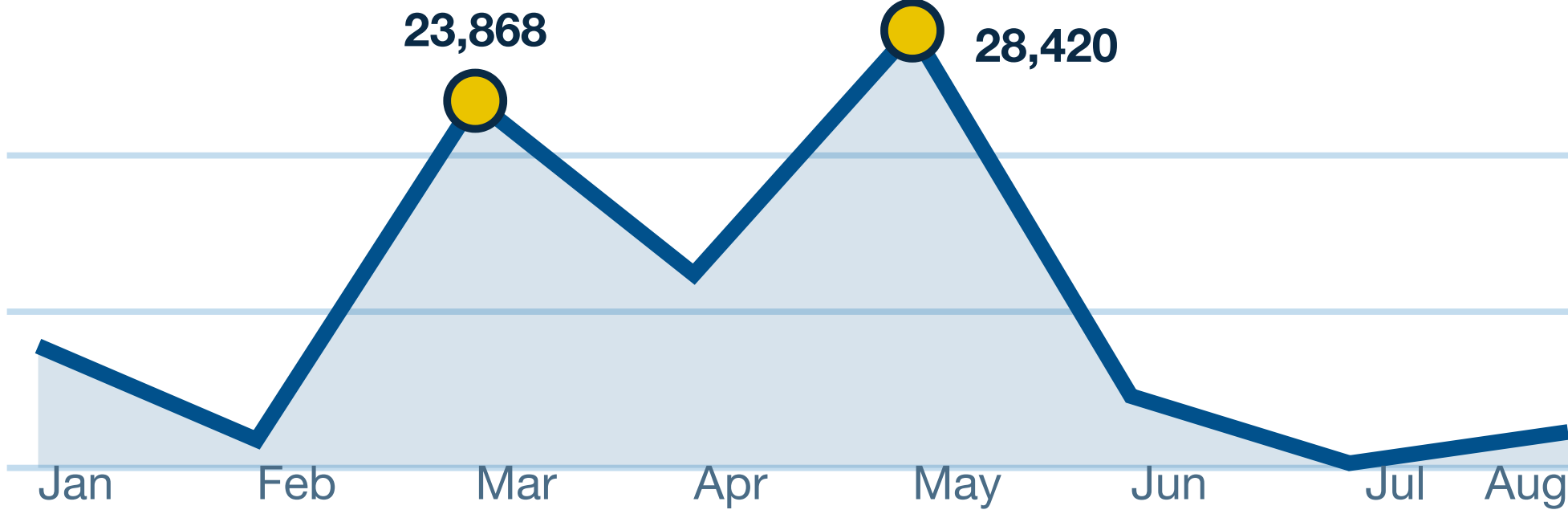


RESULTS · IN PRODUCTION



RESULTS · RHYTHM

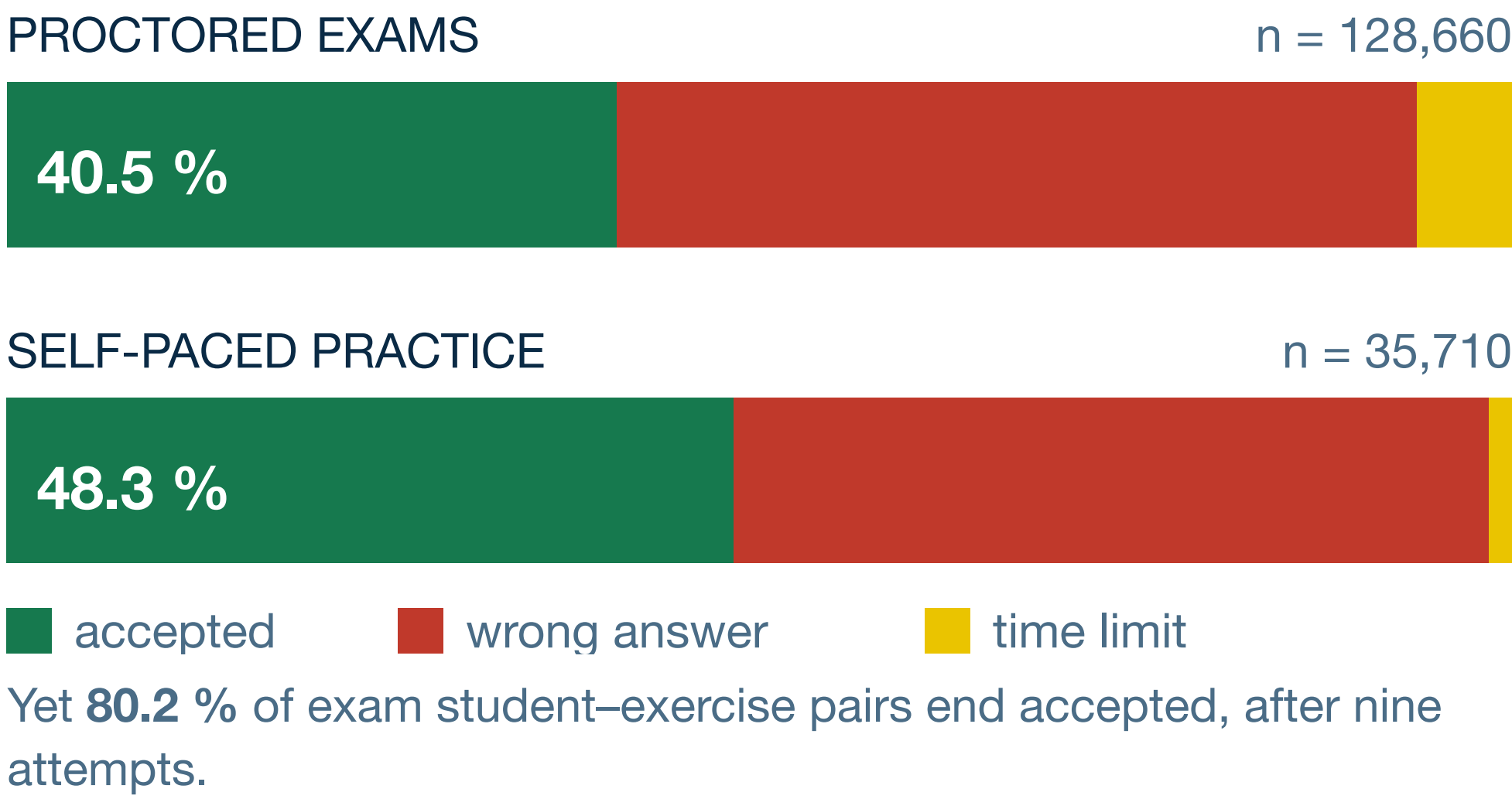
Peaks are exam weeks



Exam submissions per month, 2026 — 81,989 so far, against 46,639 in all of 2025.

RESULTS · OUTCOMES

Exams are measurably harder



RESULTS · ASSISTANT PILOT

163 conversations, 31 students

- 6 turns asked it to solve the task — all refused, all logged.
- 4.5 s median reply on a self-hosted 14B model.
- 144 of the 163 turns fell in the last 30 days.
- Practice mode only — no learning-gain claim yet.

CONCLUSIONS

- RQ1** — levels + logs make AI help auditable.
- RQ2** — verification, not generation, is the bottleneck.
- RQ3** — a curriculum judge needs stateful sandboxes.