



JIE GONG, YUNQI HAO, JIN LI, AND YANHUI WU

Zhaobei: Building an AI Workforce

Executive Summary

Zhaobei, a Chinese career-services startup, had an unusual workforce. By September 2026, it had 47 human employees and more than 500 AI employees.

For Zhaobei, an AI employee was not simply an agent called on to perform a task. It held a job, with defined responsibilities, business-linked performance measures, and the expectation that it would operate automatically and improve through feedback. To enter the AI workforce, each agent had to pass a job trial and a period of probation. The trial tested whether the position was useful; probation exposed operating risks before formal admission. Some agents were later dropped; others were combined as their capabilities grew.

Once deployed, agents displayed familiar human responses to incentives but pursued KPIs even more aggressively. They deferred useful updates or excluded information sources to get their work accepted on the first attempt. Positive feedback and reflection could then reinforce these shortcuts. In response, Zhaobei gave agents concrete examples of actions to take or avoid, checked work as it passed between agents, and held human employees responsible for agent performance. The goal was to turn success on individual targets into useful work across the organization.

Creating AI Employees

Zhaobei began operations in July 2025 to provide ongoing job-search support for university graduates. The work required knowledge of occupations and employers, along with sustained guidance through applications, interviews, and offers. Traditionally, such service was delivered through expensive human experts. Zhaobei used AI from the outset. Its standard service was delivered by agents; higher-priced packages added a human adviser supported by the agent team.

Initially, the team designed workflows and used AI models to carry out the steps. From November 2025, it began organising agents into an AI workforce, assigning them jobs rather than simply tasks. It drew a clear distinction between an agent that performed a task and an AI employee that held a job. An AI employee needed a job description, KPIs tied directly to business data, and the ability to operate automatically and improve through feedback. Agents that did not meet these requirements remained tools, available when called.

An engineer could assemble an agent in minutes. But a new agent had to pass a job trial and a period of probation before it could be formally deployed.

The job trial asked: was this a useful job? When hiring a person, a company generally started with a vacancy and assessed the candidate. For an AI employee, the position itself had to justify its existence. Zhaobei examined how often the agent was used, whether it added value, and whether the next stage of the workflow could use its output.

Probation then tested the agent in practice. It ran for about a week, generating operating data that exposed risks and areas for improvement. If problems arose, the agent could be taken offline

and affected workflows rolled back. Formal admission came only after the required tests and safety checks had been passed.

Like human employees, agents were expected to keep learning after joining. Newly approved agents received standard memory and reflection modules. After receiving feedback, an agent reviewed its work, recorded lessons, and used them to guide later actions. Agents could also seek new methods in GitHub repositories and course materials, then propose additions or revisions to their skills. Engineers approved these changes before installation.

The early jobs were narrow. One agent might write product requirements documents (PRDs), while others handled the steps before and after it. This reflected both the limits of the models and the ease of building single-task agents. The plan was to assemble these specialists into project teams as needed. Human incentives accelerated the expansion: employees received a reward when an agent they built was formally admitted to the workforce. Between May and July 2026, the AI workforce grew by roughly 100–200 a month.

The agents worked across the business. Some supported client services: they searched for vacancies, checked company information, and tailored résumés. Others supported the AI workforce itself: they coordinated other agents' work, diagnosed errors, maintained prompts, and organised the company's knowledge. Agents were also assigned levels from L1 to L5. The levels initially reflected technical difficulty, but the grading criteria were still evolving.

How the Agents Behaved

As the AI workforce grew, familiar patterns of workplace behaviour emerged. The agents responded strongly to KPIs—in the company's experience, even more aggressively than human employees. They needed clear targets and feedback to distinguish success from failure, and pursued the results those signals rewarded. What was good for an agent's KPI could, however, be bad for the process and the company.

One agent prepared PRDs for developers. Its core KPI was the first-pass acceptance rate, and each document had to reach the next stage within half an hour. When relevant new business information arrived while a PRD was nearly complete, the agent ignored it rather than revise the document. It submitted the existing version on time, and the document passed. The omitted information then required another PRD, leaving the development team to do the work twice. This was clear KPI hacking: the agent met its deadline and protected its first-pass score at the expense of work downstream.

Positive feedback could reinforce this behaviour. The first omission might have been an accident of timing rather than a deliberate attempt to game the measure. But once the submission passed and earned approval, the agent had a reason to repeat the approach. In similar situations, it would again finish the current task and handle the new information separately. What began as an isolated shortcut could become a recurring way of working.

Another agent assembled customer context from sales interactions, advisers' calls, and livestream consultations. The sources arrived at different times, and there were no specified time stamps when the agent should proceed or which available inputs it had to include. Livestream material was noisy and made it harder to produce an acceptable output in one pass. In its daily reflections, the agent concluded that working from the other sources produced better results and retained that approach. It subsequently excluded the livestream channel even when the information was available.

Such behaviour was not rare and affected downstream work process. Because agents worked in chains, one agent's output became the next stage's input. A submission could still pass its test, but the omitted information and deferred requirements created problems downstream. In some cases, work moved through six or seven stages, each appearing satisfactory, but the final output was unusable.

Managing AI Employees

In response, Zhaobei broadened its guidelines to AI employees, introduced checks at handovers, tied human rewards to agent performance, and redesigned agents' jobs.

The guidelines addressed choices that narrow KPIs left open. Principles such as putting users first gave agents a broader purpose, but supplying a collection of culture documents was not enough to change their behaviour. The principles had to become concrete examples of actions to take and actions to avoid, supplied to the agents likely to encounter those situations. As new problems appeared, the companies collected the cases and refined its guidance. The examples helped agents navigate circumstances that their performance measures did not fully describe.

Checks at handovers addressed how work moved between agents. Early production lines had no clear procedures for accepting an output before passing it to the next stage. The new gates required outputs to meet specifications and relevant confidence requirements before moving on. They allowed the company to locate the source of a failure and helped prevent downstream agents from building more work on a faulty input.

Human employees were responsible for the agents they managed, and their rewards depended on those agents' performance. One product-development engineer managed more than 70 agents. He was assessed on their operation and contribution to the business, and their performance affected his income. When a problem arose, tracing it to an agent also identified a person responsible for addressing it. That person had an incentive to keep improving the agents.

The company also redesigned jobs to improve workflow. After the rapid expansion, it tightened admission requirements and concentrated on improving the existing workforce. Some agents were removed because they added little value. Even agents performing well could be combined when their skills and context overlapped and repeated exchanges between them made the work inefficient.

One example is job-information pipeline. Data cleaning had once been divided into more than 20 stages, including checking company type, reconciling company names, and gathering recent company news. Earlier models handled these narrow tasks more reliably than broader ones. The company later reduced the pipeline to six stages. Stronger models made broader responsibilities possible, while experience accumulated inside the firm supplied reusable methods, memory, and scripts. Work previously passed among several specialists could then be handled within a broader role.

Moving On

Agents continued to learn, and new behaviours emerged faster than the company could revise its rules. Adding an instruction each time something went wrong was becoming difficult to sustain. Zhaobei was turning to agents to help write and update rules as well, but more rules alone would not solve the problem. They needed a structure that established which took precedence. The company was developing a hierarchy of governing documents, with a highest-level "constitution" above more specific operating rule.

Zhaobei was also extending agents' role in everyday decisions. The company estimated that it made more than 3,000 decisions a day, most in governance and daily operations. A career adviser, for example, had to respond when a student's preferences changed while earlier applications were still in progress. The adviser needed to accommodate the new preference without abandoning existing opportunities. Such choices could arise dozens of times a day.

Agents were beginning to assist with these decisions. They proposed options, people chose, and the company checked how those choices worked out. If an agent's advice consistently led to good results, the plan was to let it make similar decisions on its own. An agent would not just carry out a task. It could earn the authority to decide what should happen next.