

Team Pulse by CoastZone — The Deep Dive

Week 31, 2026 (Tuesday 28 July 2026)

Teams, leadership and behavior in late July 2026: speed is masquerading as value

Companion read to the podcast. Written for consultants working in a Danish/Nordic context. Every claim links to its source so you can dig in.

The big picture

This week's four stories share a single, slightly uncomfortable thread: **the clock is deceiving us**. Fast replies get read as competence in hiring. AI-assisted speed gets read as productivity in performance reviews. Change gets launched faster than brains can absorb it. And the slow, repetitive practice through which juniors once built judgment is being absorbed by AI. In every case, the slow parts — verification, judgment, thinking time — are quietly becoming the scarce asset.

Three of the four anchors are brand new: a [McKinsey Quarterly piece on rebuilding the expertise pipeline](#) (July 14), an [HBR argument that performance management needs new metrics](#) (July 2026), and [new Management Science research on “quick-reply bias” in hiring](#) (HBR, July 15). The fourth, [McKinsey's note on the cognitive cost of change](#) (July 2026), picks up exactly where Week 26's change-overload theme left off. Four themes shape the week.

Theme 1 — Who trains the next generation of experts?

In Week 27 we covered how AI is hollowing out entry-level work. This week brings the constructive sequel: **if the routine work that used to be the apprenticeship is gone, how do you build judgment on purpose?**

What's moving. [McKinsey Quarterly's Building expertise in the age of AI: Who trains the next generation?](#) (Bryan Hancock & Charlotte Seiler, July 14, 2026) argues the talent pipeline is not disappearing — it is being **rebuilt**, and organizations that redesign it deliberately will own the next generation of experts. Two senior Microsoft engineering leaders, Mark Russinovich and Scott Hanselman (*Communications of the ACM*, April 2026), put the mechanism bluntly: AI gives senior engineers an **“AI boost”** while imposing an **“AI drag”** on juniors who lack the judgment to steer and verify AI output. The resulting incentive — *hire seniors, automate juniors* — quietly dismantles the bottom of the talent pyramid every senior role depends on. Their prescription: keep hiring juniors, accept that they initially reduce capacity, and make their growth an explicit organizational goal.

The evidence. The labor-market backdrop, as compiled in the McKinsey piece: recent US college graduates faced roughly **5.7% unemployment** in Q1 2026, with about **four in ten underemployed** (NY Fed). Stanford Digital Economy Lab research (Brynjolfsson, Chandar & Chen) finds workers aged 22–25 in the most AI-exposed occupations have seen a **16% relative decline in employment**. An honesty note the article itself makes: how

much of this is AI remains genuinely contested — NY Fed economists estimate the rise of **remote work** (which makes it harder to train novices at a distance) accounts for much of the increase, and **Yale’s Budget Lab finds no clear economy-wide AI fingerprint yet**. Meanwhile, a March 2026 Strada Education Foundation survey of ~1,500 executives found those expecting AI to *increase* entry-level hiring outnumbered pessimists nearly **three to one** — but a third of employers also reported AI has already stripped out the foundational, skill-building tasks juniors learn from. Whether by algorithm or by distance, the informal apprenticeship can no longer be taken for granted.

The most practical idea in the piece is the **“answer-key model”**: the junior attempts the task independently first, the AI “grades” the attempt, and a manager talks the gap through. The clinical evidence behind it is striking: simply giving physicians an LLM barely improved their long-term diagnostic performance, but a workflow requiring them to **compare and reconcile their own reasoning with the AI’s** lifted future performance to the level of the AI itself (Goh et al., *JAMA Network Open* 2024 / *Nature Medicine* 2025). In another experiment, first-year medical students who ran a five-day attempt-then-feedback loop **outperformed second-year students** — a full year ahead in training — on the targeted diagnoses, with the effect persisting at a two-week follow-up (Kiyak et al., *Journal of Surgical Education*, Oct 2025). The narrowing gap between the human attempt and the AI output is itself the metric: measurable judgment formation — something apprenticeship never had.

A concrete case. Bank of America is holding its 2026 class steady at nearly **4,000 interns and campus recruits** from 500+ schools while explicitly redesigning those roles around AI from day one, using simulation to compress the judgment-building that junior bankers once accumulated through routine work. “We’ve got to give people the experiences in a simulated way quickly,” says Josh Bronstein, the bank’s head of global talent. Its chief people officer calls the approach “intentional and long term.” For structural inspiration, Brookings’ Molly Kinder suggests borrowing the **medical residency model** — protected, structured periods of deliberate skill-building with progression toward independence as an explicit commitment; Russinovich & Hanselman propose the same one level deeper with a **“preceptor” model**, where a senior formally mentors a small cohort, observing what juniors accept, reject and misjudge in AI output.

The Danish angle. Denmark’s cultural default — *mesterlære* and *sidemandsoplæring* — is precisely the informal apprenticeship now at risk. Væksthus for Ledelse’s classic *Pas på trinet!* makes the underlying point: judgment is built by *doing* the work, not reading about it. When AI takes the routine, the Danish apprenticeship instinct is an asset — but only if someone designs the reps deliberately instead of assuming osmosis still works.

What to do Monday. Run the answer-key loop with one junior this month: they solve the task independently first, compare against the AI’s output, and you talk the gap through. Track whether the gap narrows quarter over quarter — that is your new apprenticeship metric.

Theme 2 — The cognitive cost of change: agree what stops before anything starts

In Week 26 we covered MIT Sloan’s warning that employees can absorb one or two major changes a year while leaders plan three or four. This week McKinsey takes the next step: from *counting* initiatives to *budgeting the thinking capacity* they consume.

What's moving. McKinsey's *Designing for the cognitive cost of change* (July 2026) treats **cognitive capacity — the mental bandwidth people need for deep thinking and learning — as an organizational resource to be managed strategically**, on par with budget and headcount.

The evidence. The article draws on **nine research sources** from journals including *Organization Science* and the *Journal of Management Studies*, and outlines **three practices for safeguarding cognitive capacity**. The most concrete: **agree explicitly on what work stops before something new starts**. As HR analyst Brian Heger — who has built **practical one-pagers** on exactly this — frames it, the quiet killer is “**goal creep**”: new goals and work added throughout the year without deciding what makes room for them. The problem is rarely any single change; it is the cumulative load landing on the same brains.

A concrete case. Heger's own mid-year recalibration tool is the applied version: a structured mid-year review where goals, priorities and capacity are explicitly re-balanced — with removal decisions given the same formality as launch decisions.

The Danish angle. Dennis Nørmark's *pseudoarbejde* is the Danish shortcut into this theme: much of what should stop is already work without value. And Denmark's trust culture is an advantage here — it is culturally easier to say out loud “we are shutting this down” in a flat, low-power-distance organization. But someone has to go first, and that someone is the leader.

What to do Monday. Institute the “stop meeting”: before the next initiative launches, the team writes down — in writing, not in spirit — what stops to make room. No stop list, no new project.

Theme 3 — The performance paradox: your KPIs reward AI speed and punish judgment

What's moving. HBR's *Performance Management Needs New Metrics in the AI Era* (July 2026) names something many leaders sense but haven't formalized: organizations still evaluate people with pre-AI metrics like productivity and goal completion — creating a **performance paradox**. The employee who leans heavily on AI looks highly productive; the employee who slows down to verify and correct AI output looks less efficient — **even though that is often where the most value is added**.

The evidence. A field experiment published in *Organization Science* put 750 knowledge workers on tasks with GPT-4. Within the AI's capability, workers were **25% faster and 12% more likely to succeed**, with higher-quality output. But on a task just *outside* the AI's capability, AI users were **19% less likely to reach a correct solution** than colleagues working without it. Speed and correctness diverge exactly where the AI's competence ends — and your metrics can't see the boundary.

The framework. The HBR authors propose three measurement layers: metrics for the **human contribution**, metrics for the **AI system**, and metrics for **how well the two perform together** — including a “**complementarity index**”: did the person genuinely add value beyond what the AI produced, such as catching an error or reframing a problem? (Fair caveat, as Brian Heger notes: the index is conceptually clear but still needs operationalizing — who judges “incremental value”, and at what standard?)

The Danish angle. This is one place the Danish leadership tradition has a head start. Complementarity is best judged in **dialogue, not dashboards** — and **trust-based leadership** with close, honest 1:1s is exactly the format where “where did you save the AI from being wrong?” becomes a natural question rather than surveillance.

What to do Monday. Add one question to your next 1:1: “*When did you last catch an error, or reframe a problem, that the AI didn’t see?*” That answer is the value your current KPIs cannot measure — start collecting it before you redesign any scorecard.

Theme 4 — Behavioral science: quick-reply bias, when response time beats qualifications

What’s moving. New research in *Management Science*, presented in HBR as *Are You Biased Toward Job Candidates Who Reply Quickly?* (Eric M. VanEpps, Vanderbilt & Einav Hart, George Mason; July 15, 2026), documents a hiring cue almost nobody admits to using: **how fast candidates reply to messages powerfully influences who gets hired.**

The evidence. An analysis of **more than 11 million transactions on Fiverr** plus a series of controlled experiments found that **even a one-hour delay can sharply reduce a candidate’s chances — sometimes outweighing higher ratings or stronger applications.** The mechanism: decision-makers implicitly read fast replies as signals of **competence, warmth and future responsiveness** — yet few recognize how much they rely on the cue. The authors’ advice to leaders: explicitly decide *when speed truly matters* for the role, so you stop overlooking strong talent; and to job seekers: respond quickly, but authentically.

A concrete case. The Fiverr dataset is the case: a real marketplace where hiring decisions at scale tilted toward the fastest repliers, not the best-rated providers.

The Danish angle. This is the anatomy of *mavefornemmelsen*. Lederweb’s *Derfor føles det rigtigt at ansætte den forkerte* points to the same remedy: structure the assessment so bias doesn’t get the casting vote — especially in Denmark’s informal hiring culture, where “god kemi” often carries heavy, unexamined weight.

What to do Monday. Before your next hire or vendor selection, write down *in advance* whether response speed actually predicts success in the role. If it doesn’t, don’t let the inbox vote.

Books & tools to watch

- **Matt Beane** — *The Skill Code* (Harper Business, 2024 — backlist, but the intellectual backbone of this week’s Theme 1): skills develop through challenge, complexity and connection to expert thinking — the exact conditions AI erodes.
 - **Molly Kinder (Brookings, Jan 2026)** — the **medical-residency argument for entry-level roles**: worth citing in client conversations about early-career design.
 - **Brian Heger’s free one-pagers** on **cumulative change load and mid-year goal recalibration** — simple, client-ready versions of Theme 2.
-

The three things to remember

1. **The pipeline isn't disappearing — it's being rebuilt.** Design the apprenticeship deliberately: answer-key loops, simulation, preceptor-style mentoring. The narrowing gap between the junior's attempt and the AI's output is your new learning metric.
2. **Cognitive capacity is an organizational resource.** Agree on what stops before anything new starts — the price of change is paid in the team's thinking time, not just in the project plan.
3. **Measure complementarity, not just speed.** Reward the person who catches the AI's error — and check whether "fast" is quietly deciding your hires and evaluations.

Where in your organization is "fast" currently being mistaken for "good" — and who is paying for it?

Sources & dig deeper

International research & media - McKinsey Quarterly — Hancock & Seiler, [Building expertise in the age of AI: Who trains the next generation?](#) (July 14, 2026) - McKinsey — [Designing for the cognitive cost of change](#) (July 2026); context via [Brian Heger](#) - Harvard Business Review — [Performance Management Needs New Metrics in the AI Era](#) (July 2026); VanEpps & Hart, [Are You Biased Toward Job Candidates Who Reply Quickly?](#) (July 15, 2026) - *Organization Science* — [field experiment: 750 knowledge workers using GPT-4](#) - Brian Heger / Talent Edge Weekly — [Building expertise \(context\)](#); [Performance metrics \(context\)](#)

Danish sources & voices - Væksthus for Ledelse — [Pas på trinet!](#); [Tillidsbaseret ledelse i praksis](#) - Lederweb — [Derfor føles det rigtigt at ansætte den forkerte](#) - Dennis Nørmark — [Pseudoarbejde](#)