

MEETING SUMMARY

Adopting AI and New Technology in Small and Midsize CPA Firms

Date: August 10, 2026 | **Speaker:** Dr. Lee Rogers, CPA Crossings

■ Quick Recap

Dr. Lee Rogers presented findings from the publication *Start From Trust*, arguing that the central question for small and midsize firms is not whether to adopt AI but whether that adoption is governed or unstructured. Empirical research shows people abandon an AI tool after a single mistake while forgiving the same error from a colleague, which makes calibrated trust — matching confidence to the tool's demonstrated capability on a specific task — the practical starting point. Rogers introduced a two-axis framework for assessing a firm's posture, examined why AI tends to empower newer staff while threatening the professional identity of experienced ones, reviewed research on when human-AI pairing helps and when it degrades results, and closed with five questions to put to any AI vendor and five ways for a firm to begin engaging deliberately.

■ Summary

The Trust Problem

Research shows an asymmetry in how people extend trust: a single visible mistake is often enough for someone to abandon an AI model outright, while comparable human error is absorbed and forgiven. That asymmetry shapes adoption across the profession. Sixty-three percent of CPA firm managers report a "wait and see" posture, and while concerns about privacy, ethics, and accuracy are widely voiced, actual deployment of generative AI sits at roughly six percent. Firms are, on the whole, choosing careful consideration over rapid movement.

Calibrating Trust, and the Four Postures

Drawing on a 2004 article by engineering psychologists on human trust in machines, Rogers argued that trust should be matched to a tool's capability for a particular task rather than granted or withheld wholesale. He set out five approaches ranging from reckless adoption at one end to informed caution and responsible piloting at the other, with the goal of locating "defensible ground" in the middle band. He then introduced a framework built on two axes, trust and usage, producing four postures for putting AI into a workflow. None of the four quadrants is wrong in itself; the value of the exercise is in showing where an organization needs more structure and where it may need a new initiative.

Identity Threat Among Experienced Professionals

Research on how experienced professionals react to AI tools identifies two distinct identity threats: doubt about one's own capability to perform a task, and doubt about continuing to be recognized as an expert. The pattern is uneven across a firm. AI adoption tends to empower newer employees, who gain reach and speed, while unsettling those whose standing rests on accumulated expertise. Rogers was clear that this is an ordinary human reaction rather than a symptom of any one organization's culture, which matters for how firm leaders frame the conversation internally.

Unstructured or Governed: The Real Choice

Eighty percent of workers already use personal AI tools at work. That figure reframes the decision in front of firm leadership: the binary of embracing or rejecting AI has already been settled in practice, and what remains is a choice between unstructured use and governed use. Rogers offered two questions that reliably predict whether trust in a tool is warranted — whether the AI is demonstrably more capable than a person at the specific task in question, and whether that task requires individualized judgment.

When Human-AI Pairing Helps, and When It Doesn't

Combinations of human and AI generally outperform humans working alone, but the finding comes with an important qualification: where one partner is materially stronger at a task, adding the weaker partner reduces performance rather than improving it. Rogers presented two studies pairing physicians with GPT-4. The AI performed well on closed diagnosis tasks with determinate answers, and underperformed human experts on open judgment tasks calling for nuanced decisions. He also warned of the "explanation effect" — people accept an AI recommendation more readily when it arrives with an explanation, whether or not that explanation is sound.

Five Questions to Ask AI Vendors

Rogers set out five trust questions to put to any vendor: how permission structures work, how the system handles uncertainty, what audit trails exist, how outputs can be explained, and what change control procedures are in place. He argued that trust and speed are better understood as nested rings than as opposing values, with the technology decision sitting at the center, encircled by regulatory duty, professional judgment, and client relationships. He also noted a tension worth watching: firms are currently seeing financial growth, but that growth can mask underlying vulnerabilities in talent pressure, shifting client expectations, and deferred learning.

Cognitive Offloading and the Talent Pipeline

Eighty-eight percent of firms report concern about rising client expectations, eighty-seven percent about talent shortages, and eighty-five percent about AI replacing entry-level roles within two years. Against that backdrop Rogers described cognitive offloading, where skills atrophy once a tool takes over work that previously required sustained mental effort. Offloading some skills to technology is ordinarily acceptable, he noted; the real concern is how new professionals will develop judgment and critical thinking when the nature of entry-level work changes this quickly. He urged skepticism toward sweeping claims that Google has ruined memory or that ChatGPT is reshaping psychology, while acknowledging that technology adoption does produce measurable changes in brain structure.

Five Ways to Engage

Rogers closed with five practical steps: inventory the technology the firm already has, run low-stakes pilots, establish temporary rules rather than permanent policy, form pilot teams, and set review dates in advance. Alongside those, he offered guardrails for day-to-day use — give human reviewers genuine control over AI output, require independent judgment before the AI's assessment is seen, make the tool available for interrogation rather than accepting it at face value, and stay mindful of how confidently a recommendation is expressed