

Breakout Questions

1. **Who do you trust?** *The people running your OT? The AI running on your systems? The AI running your systems and hopefully making the right decisions? Your suppliers and their supply chains?*

Prompts to Consider:

- What's changed for your 'circle of trust' in the past 2 years?
- Do you consider internal systems in the same way as external ones? If not, why not?
- What is your current evidence that your systems are doing what you think they are? Is that evidence sufficient?
- What level of architectural accountability would be sufficient — both for approving AI workloads and for resisting AI-enabled attacks?
- How will you know what kind of software you are dealing with? (e.g. if AI designed the drivers for your software)
- What do you require in a trusted supply chain?

2. **Who is your adversary?** *AI has made your adversaries more capable and your systems harder to change with confidence. Where are those two pressures hitting you hardest?*

Prompts to Consider:

- Is it time to reframe who your adversary actually is? Previously it was nation-state actors. Now it could be anyone from organised crime to someone down the street. Is your current security posture keeping pace?
- What decisions are you deferring (updates, AI deployments, architectural changes) because the cost or risk of change is too high?
- Where are those two pressures compounding each other i.e. systems that are increasingly exposed and increasingly difficult to update?

3. **What does AI mean in OT?** *Can your frameworks keep pace? Traditional OT frameworks were built for a different threat environment and a slower rate of change. If you could update your OT systems with confidence in days rather than months, what would you do differently?*

Prompts to Consider:

- How will you deal with the rate of update churn? Traditional OT frameworks were designed for updating once a year, now you might get three a week.
- How should OT regulators deal with this new normal?
- The frameworks were designed for a different threat environment and a different rate of change. Which are most urgently broken — and who moves first to fix them: industry, regulator, or vendor?

- What response timelines would feel acceptable for security patches?
- What does the competitive and mission picture look like if your certification cycle compresses from months to days — not just for security, but for operational agility and the ability to deploy new AI capabilities when the mission demands it?

4. *When the AI on your system fails...*, behaves unexpectedly, or is compromised
– *what does your architecture actually do?*

Prompts Consider:

- Do you have a mechanism to detect AI failure in real time? To isolate it, remove it, or recover from it without shutting down operations?
- How much of that response is automated, and how much depends on a human being in the right place at the right time?
- And if you don't have clear answers to those questions, what does that tell you about your readiness to deploy AI at scale on critical systems?