



## **AI in Cocoa: Separating Signal from Noise for Senior Leaders**

### **Introduction**

Building on the themes raised in "The Bittersweet Byte," this session moves from diagnosis to action. You'll learn how the sharpest business minds stress test a plan before committing to it, then apply that same rigor to AI, so you can judge whether an opportunity is worth pursuing and get useful, reliable output from AI tools day to day.

### **Who should attend?**

Senior Leaders (EVP, VP, Directors) and aspiring leaders, and decision-makers across the cocoa and chocolate value chain, traders, processors, manufacturers, and cooperative leaders, who need both a framework for evaluating AI investments and practical skills to use AI tools well themselves.

### **Expected Takeaways**

- How top operators pressure test decisions before they commit: Bill Gates' habit of hunting for the flaw that kills an idea, Elon Musk's first principles approach of stripping a plan down to what's actually true, Warren Buffett's discipline of staying inside your circle of competence and asking what would have to be true for this to fail, and the McKinsey partner's habit of forming a hypothesis first and then trying to disprove it
- How to apply that same discipline to an AI plan or pilot before it launches, not after budget is spent
- How to tell a solid AI investment from a feature list dressed up as a solution, drawing on what has worked (and failed) in cocoa and comparable crops like rice, coffee, and palm oil
- What business viability, data readiness, and farmer impact actually look like in practice
- How to write clear, effective prompts: giving AI the right context, framing the task, and specifying the output you actually want
- How to ask AI for analysis, comparison, or summarization of things like market reports, supplier data, or sustainability disclosures

## Business outcome:

Leaders leave with three things: the mental models the best operators use to flip a plan on its head before it proceeds, a framework for deciding where AI is worth real investment, and the practical skill to prompt, question, and validate AI outputs themselves, rather than taking any team's claims (vendor or internal) at face value.

- How to validate AI outputs and avoid blind trust, whether the claim comes from a vendor pitch or your own team's pilot results, especially with figures, sourcing claims, or anything that touches compliance
- A simple, repeatable filter for evaluating the next AI pitch or tool that lands on your desk, regardless of where it comes from

## Speaker Bio



William Gaultier is a globally recognized advisor to Global 500 boards and UHNW individuals, operating at the high-stakes intersection of **Agentic AI, cybersecurity, and operational reputation**. A seasoned founder and operator, William has built and led four companies and held VP-level roles at **Sephora** and **Lazada**, where he managed \$200M+ P&Ls and generated over \$2B in cumulative revenue.

As a Partner at **Enfactum**, William specializes in neutralizing the AI "Actionability Gap". He pioneered the **HP Garage 2.0** innovation hub in APAC, curating a vibrant ecosystem of AI startups and global partners like Intel and Nvidia to drive edge-AI commercialization. His hands-on entrepreneurial background is balanced by deep regulatory expertise; he has advised US and APAC governments on cross-border compliance and directed Sephora's 8-country response to a major data breach.

Today, William equips leadership teams to treat modern risk as existential. By applying Enfactum's **5-Step AI Growth Assessment**, he helps boards design resilience into their people and processes, ensuring that reputation moves as fast as code while maintaining strict human-in-the-loop governance.