

RISKBOWL LIVE: MODEL RISK & AI ROUNDTABLE

On 9th October 2025 we held our latest RiskBowl Live roundtable with participants from 10 banks and building societies, alongside two of our senior advisors: Colin Jennings (ex-PRA and ex-CRO) and Lukasz Szpruch (The Alan Turing Institute).

This roundtable brought together senior heads of Model Risk Management to take stock of where banks are on managing the model risk of AI and discuss their convergence towards full compliance with SS1/23, under Chatham House rules.

The discussion confirmed a common trajectory: an early phase of AI modelling experimentation has exposed structural gaps — in taxonomy, inventory management, model monitoring and validation — that now require a coordinated effort to make bank-wide AI use safe, auditable and scalable across firms. Firms have welcomed the clarity and the heightened stature of Model Risk in the firm's risk taxonomy, and are using its guiding principles to coordinate said effort.

Key takeaways from the discussion are presented below:

Managing the model risk of AI

- Experimentation to consolidation: participants described an early period of numerous disjointed pilots and recommend grouping similar use cases to scale efficiently rather than proliferate ad-hoc AI projects
- Use-case-specific governance: high-risk algorithmic/decisioning use cases require materially different controls from low-risk productivity tools; a one-size governance model is insufficient
- New tech stack & MRM implications: generative AI brings dependencies that must be formally approved and governed. These create integration and approval work that traditional MRM processes were not designed to cover.
- Skills and vendor risk gaps: many teams or vendors originate outside banking and lack knowledge of bank's processes, compliance expectations, and model lifecycle controls; stronger third-party standards and onboarding are needed
- Model classification ambiguity: simple assistive tools (e.g., grammar correction) may fall outside current model definitions, while some AI systems sit partially within model risk remit—creating uncertainty about monitoring and ownership
- Committee and oversight design: avoid duplication of oversight bodies — firms must clarify roles between existing model risk committees and any AI monitoring forums
- Quantitative monitoring and human-AI controls: firms want monitoring frameworks capturing model *and* human performance, with defined escalation triggers and the ability to switch to automated testing based on scale and level of risk

Convergence towards compliance with SS1/23

- Raised standards and visibility: SS1/23 has driven broader Model Risk visibility within firms and heightened board awareness
- Material operational uplift: documenting and managing additional models and DQMs in scope is increasing resourcing and cost materially — firms reported significant headcount increase and process redevelopment
- Definition and scope tensions: debate continues on what counts as a model (quantitative, deterministic, qualitative outputs, agentic behaviours) and on incentives to classify or de-classify to manage control and operational burden
- Validation and ownership challenges: validating qualitative and AI-enabled outputs is resource intensive; teams need clarity on who conducts testing (first line, MRM, or specialist validation units) and on practical monitoring cadences
- Ongoing dialogue required: participants agreed continued cross-firm engagement and proactive regulatory conversations are necessary to align interpretations and reduce operational fragmentation between firms and subsidiaries



Cem Dedeaga

Partner, Head of Risk Modelling UK&I

cem.dedeaga@oliverwyman.com



Matias Coggiola

Senior Manager, MRM lead

matias.coggiola@oliverwyman.com
