

# Case Study: “RiskMind AI” at GlobalRisk Brokers

## Background:

GlobalRisk Brokers is a leading insurance & reinsurance brokerage firm serving multinational insurers and cedents. They’re launching **RiskMind AI**, a Gen AI–powered platform that analyzes historical claims, market data and treaty terms to recommend optimal reinsurance structures and pricing. To ensure responsible, compliant, and trustworthy deployment, they’re implementing the four governance mechanisms from the Responsible AI mind-map.

---

## 1. Regulatory Sandboxing

- **Scoped Pilot Environment**
    - **Legal Test Zone:** Partnering with the UK’s FCA sandbox to test RiskMind AI on real cedent portfolios under relaxed reporting deadlines.
    - **Entry/Exit Criteria:** Must hit  $\geq 90\%$  pricing accuracy &  $\leq 2\%$  fairness drift before full rollout.
  - **Regulator–Developer War Room**
    - Weekly working-group calls with FCA analysts to surface emergent model risks (e.g., bias against certain geographies).
  - **Metric-Driven De-Risking Plan**
    - Track false-positive rate (e.g., over-priced treaties), throughput (processing time), and explainability scores as go/no-go gates.
  - **Adaptive Policy Iteration**
    - Use sandbox learnings to refine both internal controls and help shape FCA guidance on AI in reinsurance.
- 

## 2. AI Assurance Frameworks

- **Layered Assurance Taxonomy**
  - **Tier 1 (Low Risk):** Advisory modules on historical loss trends—self-certified quarterly.
  - **Tier 2 (High Risk):** Automated treaty-pricing recommendations—mandatory third-party audit per EU AI Act Annex IV.
- **Continuous Assurance Pipelines**
  - CI/CD gates that run bias-tests on cedent segments, robustness checks under stress scenarios, privacy scans over PII, and security scans for model APIs.
- **Assurance Labels & Trust Marks**
  - Issue an internal “AI-Trusted” badge for each module that passes all gates—visible in the RiskMind UI.

- **Hybrid Audit Models**
    - Low-risk components use in-house reviews; high-risk treaty-structuring models undergo external assurance by a Lloyd's-accredited auditor.
- 

### 3. AI Board Risk Templates & Escalation Triggers

- **Standardized Risk Dashboard**
    - Board pack includes: model inventory, risk classification (e.g., pricing vs. advisory), control status, audit findings & residual risk scores.
  - **Automated Escalation Rules**
    1. **Bias-Drift** > 3 % → Ops team auto-alert
    2. **Throughput** < SLAs → Risk committee review
    3. **Critical security finding** → Immediate board notification
  - **RACI-Style Role Mapping**
    - Model Owner (Head of Analytics), CDO, CRO, Legal & Compliance each have defined sign-off at every escalation tier.
  - **Time-Bound Remediation SLAs**
    - 24 h to patch fairness deviations; 48 h to investigate any data-integrity alerts.
- 

### 4. Privacy-by-Design & Data Embassies

- **Built-In Privacy Engineering**
    - All cedent data is ingested with differential-privacy noise on small portfolios, PII is tokenized, and model outputs are encrypted at rest.
  - **Data Embassies (Sovereign Hubs)**
    - Critical client data resides in a Swiss “data embassy” under Swiss privacy law; RiskMind AI algorithms are containerized and shipped to the embassy—no raw data ever leaves.
  - **Compute-to-Data Paradigm**
    - Underlying AI jobs run within the embassy; only aggregated insights exit through audited APIs.
  - **Tamper-Proof Trust Anchors**
    - Model-version hashes and data-lineage logs are anchored in a consortium blockchain shared with key cedents and reinsurers.
-

# Participant Exercises

## A. Discussion Questions

1. **Sandbox Design:**
  - How would you refine the entry/exit criteria to balance innovation speed with risk controls?
2. **Assurance Strategy:**
  - Which model components should escalate to third-party audits, and why?
3. **Escalation Triggers:**
  - Propose an additional KPI breach trigger for the risk dashboard.
4. **Data Embassy Trade-Offs:**
  - What challenges might arise from the compute-to-data approach, and how can you mitigate them?

## B. Key Deliverables

1. **Draft FCA Sandbox Agreement**
  - Outline scope, KPIs, and learnings-to-policy feedback loops.
2. **AI Assurance Pipeline Diagram**
  - Show CI/CD gates, test suites, and badge issuance logic.
3. **Board Risk Dashboard Mock-Up**
  - Include escalation rules and RACI annotations.
4. **Privacy-by-Design Policy & Embassy Architecture**
  - Detail data flow, encryption/tokenization steps, and trust-anchor mechanisms.

Map each deliverable back to the corresponding branch of the Responsible AI mind-map and be ready to present your rationale.