

Hiring for AI in BFSI: four roles to hire & *three to skip*.

AI in financial services is shaped by regulators, not technologists. The roles that scale a bank's AI program, and the ones that look essential but quietly stall it.

Model Risk · AI

AI Governance

Fairness Audit

MLOps · Regulated

GenAI Governance

SR 11-7 · OCC · EU AI Act

Four roles to *hire now*.

Model Risk Manager (AI-specialized)

Validates and governs every model touching a customer or the balance sheet. **Red flag:** has only ever validated linear / logistic models.

Fairness Audit Engineer

Builds and runs the testing that surfaces fairness drift; sets the thresholds that trigger retraining. Its own discipline, not a 20% side-task. Tooling: Aequitas, Fairlearn, AIF360.

AI Governance Lead

Owns the policy framework, model inventory, and the regulator relationship: process work, distinct from validation. **Proof:** has shepherded a model through a federal review.

MLOps Lead (regulated environment)

Deployment, monitoring, and rollback that survive a change-control board, where you can't roll forward and watch. **Non-negotiable:** alerts meaningful yet quiet enough not to be ignored.

Three roles that *rarely deliver*.

- ✗ **Chief AI Officer.** Three jobs in one: sponsor, technical lead, thought leader. Works only with real P&L authority (an AI-first product line), not as a horizontal coordinator. Otherwise: a strong CTO plus an empowered VP of Data & AI.
- ✗ **Generative AI Strategist.** "Identify GenAI use cases" is a six-week exercise, not a standing role. Fund two senior PMs embedded in business lines to ship instead.
- ✗ **Standalone AI Ethicist.** The work is real; the isolated seat is invited late and ignored. Distribute it across Model Risk, Governance, and Fairness, the roles with implementation authority.

The question that *sorts the field*.

"Tell me about an AI model you killed. What surfaced the problem, who made the call, and what did you do next?"

Killed on fairness-audit results → understands the regulatory environment. Killed on model-risk findings → has worked inside real governance. Never killed one → never shipped, never been audited, or never had authority to escalate.

Takeaways for *your hiring plan*.

- **Audit your reqs against the four-role framework.** Hiring data scientists faster than Model Risk, Governance, or Fairness roles means the program stalls in 2026–27.
- **Question every open Chief AI Officer / GenAI Strategist req.** The same budget usually funds two embedded delivery hires that actually ship.
- **Build interview loops that test regulatory maturity,** not just technical depth. The kill-a-model question is where to start.

Related BFSI *insights*.

BFSI · RISK	FCRM, GRC, and the modern risk stack: hiring implications for 2026.	Read →
WORKFORCE	Direct hire vs. contingent: a procurement leader's decision framework.	Read →
AI & SOURCING	AI-augmented sourcing: separating real productivity from buzzword theater.	Read →