

Decoding The RBI eMRM Guidelines

What the June 2026 Draft Requires, and Why It Matters

What the RBI's June 2026 draft requires of every regulated entity across the full model lifecycle, from board-owned governance and independent validation to a dedicated regime for AI, GenAI and Agentic Systems, and how institutions can build ahead of the final circular.



Executive Summary: What the RBI's 2026 Model-Risk Draft Actually Does

On 24 June 2026 the Reserve Bank of India issued its draft “Guidance on Regulatory Principles for Model Risk Management, 2026” (PR 2026-2027/528) for public consultation, with comments due 24 July 2026. On finalization, it will supersede the model-risk provisions of the 2002 Guidance Note and establish a single, board-approved framework governing every model a regulated entity uses to support business or decision-making.

The scope of the guidance draft is enterprise-wide, covering all models with material decision impact, statistical, business-rule, algorithmic, machine-learning, Generative-AI and third-party alike, whether developed in-house, procured, or built jointly, and applies to all eleven categories of RBI-regulated entities. It is principles-based and proportionate: control intensity is set by each model's materiality, complexity, usage and impact.

Regulated entities must act on six obligations: establish a board-approved MRM framework with control intensity proportionate to each model's risk tier, maintain a complete model inventory, tier every model by its risk profile, validate independently across the lifecycle, apply defined AI/ML controls, and retain accountability for third-party models. The sections that follow set out each obligation and the specific actions it requires.

“Core principle: a model may be outsourced, but accountability for its governance, validation and outcomes cannot. The regulated entity remains fully responsible for every third-party and AI model it deploys for its use.”

What Changed, In Three Moves

Scope: All Models, All Entities

Coverage extended from only credit models to every material model internal, third-party and AI/ML, across all eleven categories of RBI regulated entities.

Governance: Board-owned and Tiered

A board-approved MRM Framework, three lines of defence, and RMCB approval for high-risk models, with risk tiering and an anti-dilution rule.

AI/ML: Dedicated Control Set

Defined controls for explainability, bias, drift, hallucination and adversarial inputs, plus human oversight and kill-switch requirements.



The Regulatory Trajectory: The Journey to an Enterprise Model-Risk Framework

India's model-risk regime has moved from a brief 2002 credit-modelling chapter, through a decade of Basel capital-model rules, to a technology-agnostic framework covering the entire model estate, including AI, generative AI, and agentic AI models. The June 2026 draft is the culmination of that progression.

The supervisory question moves from validating only credit models to governing and evidencing the entire model estate.

Four Phases of the Regulatory Trajectory

1

Guidance Note – Chapter 3 (Oct 2002)

A descriptive credit-modelling chapter referencing the Altman Z-score and Merton/EDF approaches. It set minimum expectations only and did not address validation, model inventory or governance as distinct disciplines.

2

The Basel Capital Era (2007–2013)

Model requirements sat within the capital framework –stress-testing, ICAAP, the IMA and IRB approaches. Validation applied mainly to models feeding regulatory capital; there was no horizontal model-risk concept.

3

The Digital Lending & AI Wave (2022–2025)

Digital-lending rules, the August 2024 credit model-risk draft and the August 2025 FREE-AI Committee report extended supervisory focus toward AI-aware, enterprise-wide model governance.

4

The Enterprise MRM Framework (Jun 2026)

The draft supersedes the 2002 chapter: all eleven RE categories, all models, a board-approved MRMF, independent validation, risk tiering, non-delegable third-party accountability and a dedicated AI/ML chapter.



Seven Dimensions That Widened Under the 2026 Draft

The August 2024 draft addressed credit models; the June 2026 draft extends model risk management into an enterprise-wide mandate. The table below sets out the change on each dimension.

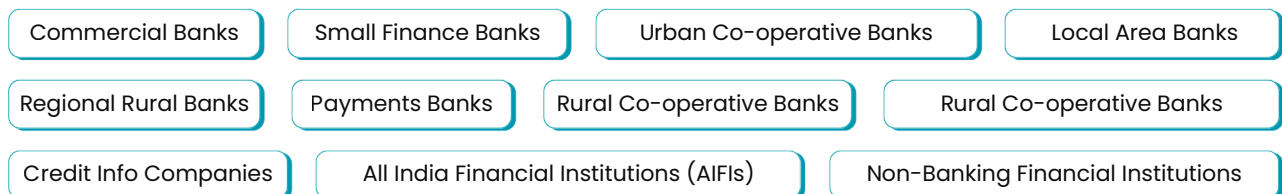
Dimension	Aug 2024	Jun 2026
Nature & Intent	A principles note scoped to credit-decision models.	A full enterprise Model Risk Management Framework (MRMF) spanning the model lifecycle.
Scope of "Model"	Credit-risk models only, scoring, rating and PD/LGD-type engines.	All models with material decision impact, non-credit, third-party, AI/ML, even rule engines.
Entities in Scope	Credit-granting REs, chiefly banks and NBFCs.	All 11 RE categories: banks, SFBs, PBs, LABs, RRBs, UCBs, RCBs, NBFCs, AIFIs, ARCs, and CICs.
Governance & Tiering	Board-approved policy, an inventory, a broad risk-based approach.	Board + RMCB approval gates, formal risk tiering with an anti-dilution rule, three lines of defence.
Third-party Models	Covered, subject to RE due diligence and validation.	Non-delegable accountability, the RE must independently validate despite any vendor certification.
AI / ML Models	Largely implicit, no dedicated regime for AI or ML.	A dedicated AI/ML chapter: explainability, bias, drift, adversarial defence, disclosure, kill-switches.
Oversight and Consumers	Limited treatment of human oversight and consumer impact.	Human-in-command mandates, override/kill-switch controls, automation-bias safeguards, consumer protection.



Current Scope: Who Must Comply, and Which Models Are Captured

The Guidance defines its reach on two axes: the entities it regulates and the models it captures. Both are drawn substantially more widely than any earlier Indian model-risk guidelines.

Applies to All Eleven Categories of RBI-Regulated Entities



And Every Model With Material Decision Impact

“Model” is defined by function, not technology, any method that processes inputs to produce an output used in a business or risk decision qualifies, whether built in-house, procured, or developed jointly. Six families fall in scope:

1. Statistical and Econometric

Regression, scorecards, time-series and PD/LGD-type models.

2. Business-rule and Algorithmic

Deterministic rule engines, limit checks and decision trees.

3. Machine-learning (ML)

Supervised, unsupervised and deep-learning predictive models.

4. Generative AI (GenAI)

LLM-based generation, summarisation and copilots.

5. Third-party / Vendor

Bought-in, foundational and externally hosted models.

6. Agentic and Autonomous

Multi-step, tool-using systems acting with limited human input.

Practical implication: Entities previously outside credit-model rules, including payments banks, ARCs and CICs, must now establish model governance from first principles across a materially larger and more heterogeneous population of models to identify, classify, validate and control.



Four Requirements

Requirement 1: Govern, Inventorize, and Tier Every Model

The Guidance requires a board-approved framework covering the full model lifecycle, a complete inventory of every model in use, and risk-based tiering that sets the depth of validation, approval and monitoring for each model.

A Board-Approved MRM Framework

BOARD & RMCB

A board-approved framework must govern the full lifecycle, development, validation, approval, deployment, monitoring and decommissioning. The board periodically reviews it, sets the risk appetite for model risk, and informs that appetite with scenario analysis and stress testing; the RMCB carries delegated oversight.

Full lifecycle

One framework spanning build to retirement, owned at the board level.

Clear Ownership

Senior management operationalises; the RMCB oversees high-risk models.

Living Policy

Reviewed periodically and as regulation and the estate evolve.

A Complete Enterprise Model Inventory

INVENTORY

A single inventory must record every model in use, active, dormant and retired—spanning statistical, ML, GenAI, third-party, and end-user tools. No model may operate outside the inventory. The register serves as the board-accountable source of truth and supports enterprise-level visibility of concentration and interdependence.

The Whole Estate

Every model and its metadata, across all lifecycle states.

Model Vs Non-model

A defensible test for what qualifies as a model on entry.

Aggregate View

Exposure, dependency and concentration made visible.

Risk-Based Tiering with Anti-Dilution

TIERING

Each model is classified by materiality, complexity, usage and impact, and the tier sets the depth and cadence of validation, approval and monitoring. Classifications must be reviewed at least annually, or on material change. An anti-dilution rule prevents a low sub-score from reducing the tier of a material model.

Tier Drives Rigour

Validation and monitoring intensity cascade from the tier.

Annual Review

Re-tiered yearly, or on any material change to the model.

No Quiet Downgrades

Anti-dilution stops materiality being used as a shield.

Requirement 2: Validate Independently + Monitor Continuously

Models must be validated independently of their developers across the lifecycle, high-risk models must be approved by the RMCB before deployment, and performance must be monitored continuously with revalidation triggered by breaches and material change.

Independent Model Validation

SECOND LINE

Validation must be independent of model development and performed before deployment, then periodically & on trigger events thereafter. It tests conceptual soundness, data quality, assumptions, methodology and interpretability & analyzes real-world outcomes through benchmarking, back-testing, & sensitivity analysis, with depth proportionate to the model's tier.

Conceptual soundness

Data, assumptions, method and interpretability, tested by tier.

Outcomes analysis

Benchmarking, back-testing and robustness against real results.

Evidenced record

A defensible validation report retained for supervisory review.

Approval & Exception Governance

APPROVAL

High-risk models require committee approval before deployment, with the RMCB reviewing their validation reports. Any model used under an exception must carry compensating controls, an accountable owner, enhanced monitoring and a documented remediation plan tracked to closure, with the rationale on record.

Approval by Tier

High-risk models escalate to the RMCB before go-live.

Exceptions Governed

Compensating controls, an owner and a remediation deadline.

Tracked to Closure

The rationale and conditions stay on the record throughout.

Continuous Monitoring & Change Control

LIFECYCLE

Monitoring must run continuously across the lifecycle: performance, stability and drift are tracked against defined limits, and outcomes reconciled with real-world observations. A limit breach must generate a tracked finding that triggers escalation or revalidation, and every model change must be version-controlled with impact-based revalidation.

Performance & drift

Metrics and breaches watched against defined thresholds.

Breach to finding

A limit breach escalates and writes back to the inventory.

Governed change

Versioning and re-validation triggered by material change.

Requirement 3: Make AI Explainable, Fair and Reliable

The Guidance sets specific expectations for AI and ML models: material decisions must be explainable, outputs must be tested for bias across affected segments, and generative systems must be safeguarded against hallucination and monitored for drift.

Explainability & Fair Treatment

TRANSPARENCY

Institutions must be able to explain how a model reaches material decisions, particularly those affecting customers such as credit, fraud or onboarding outcomes. Where full explainability is not feasible, enhanced monitoring, additional validation and compensating controls apply. Fairness and bias testing across affected customer segments is mandatory.

Explainable by Design

Interpretable methods and post-hoc techniques appropriate to use.

Fairness Tested

Disparate-impact testing across affected customer segments.

Controls Where Opaque

Extra monitoring and validation when explainability is limited.

Hallucination & Output Reliability (GenAI)

RELIABILITY

For generative systems, institutions must implement explicit safeguards against hallucination and continuously monitor output reliability, drift and uncertainty. Output quality must be measured and evidenced across systems in production, rather than relied upon on the basis of vendor assurances.

Guard the Output

Explicit safeguards against hallucination and unreliable answers.

Monitor Drift

Behavioural and semantic shift tracked as models age.

Measured, Not Assumed

Reliability quantified continuously across GenAI in production.

Supervisory expectation: explainability and fair-treatment evidence must be available on demand, and generative-AI reliability must be demonstrated through continuous measurement rather than asserted.

Requirement 4: Keep Humans in Command of AI

For AI-driven decisions, the Guidance requires human oversight with override and kill-switch controls, adversarial testing and security assessment before deployment, customer disclosure of AI use, and full retention of accountability for third-party models.

Human Oversight & Kill-Switch

IN COMMAND

High-impact AI requires human-in-the-loop or human-on-the-loop oversight, override and suspension controls, and a kill-switch to halt a system producing harmful outputs. Reviewers must have the expertise to challenge AI outputs, and explicit safeguards must counter automation bias, over-reliance and decision fatigue.

Override & Suspend

A human can intervene, override and stop the system.

Kill-switch

A hard stop for AI producing harmful or unsafe outputs.

Against Automation Bias

Reviewers equipped and required to challenge outputs.

Red-Teaming, AI Security & Disclosure

SECURITY

Customer-facing and generative systems need adversarial testing and AI-security assessment, protection against prompt injection and adversarial inputs, limits on session and context persistence, and detection of anomalous usage. Institutions must disclose to customers when they are interacting with AI, state its limitations, and offer a human alternative.

Adversarial Testing

Prompt-injection, jailbreak and edge-case red-teaming pre-production.

Secured Deployment

Session limits, anomaly detection, no new production vulnerabilities.

Customer Disclosure

Tell users it's AI, state limits, offer a human option.



Non-Delegable Third-Party Accountability

ACCOUNTABILITY

Outsourcing a model does not transfer accountability for its outcomes. The RE must independently validate vendor and foundational models regardless of certification, conduct due diligence before adoption, secure technical documentation and audit rights, and include continuity and exit provisions in every contract.

Validate In-house

Independent validation despite any vendor certification.

Contract for Control

Audit rights, documentation and exit terms secured up front.

Ownership Stays

The RE remains fully accountable for third-party outcomes.

Global Context: How the RBI Draft Compare to Other Global Standards

The RBI draft aligns with the US, UK, EU and Basel standards on the prudential core —centralized model inventory, independent validation and three lines of defence—and sets more explicit requirements than most on AI, generative AI and consumer protection.

GOVERNANCE THEME	RBI 2026	US SR 11-7	US SR 26-02	UK SS1/23	EU AI Act	Basel/ECB
Model Inventory & Tiering	✓	✓	✓	✓	●	✓
Independent Validation	✓	✓	✓	✓	●	✓
Three Lines of Defence / bBoard	✓	✓	✓	✓	—	●
Third-party Model Accountability	✓	●	✓	✓	●	●
AI / ML Governance	✓	—	●	●	✓	—
Explainability Requirements	✓	—	—	●	✓	—
Human Oversight of AI	✓	—	—	●	✓	—
GenAI / Hallucination Controls	✓	—	—	—	●	—
Consumer / Fairness Protection		—	—	●	✓	—



Full Explicit Requirement



Applicable



Partial / implied

— Not Addressed

EU AI Act (2024/1689)

Comparable AI depth, but an AI-specific regime rather than a prudential MRM framework.

US SR 26-02 (2026)

Updates SR 11-7 but excludes generative and agentic AI; the RBI brings both into scope.

UK PRA SS1/23 (2023)

The closest prudential equivalent —technology-neutral, with lighter explicit AI controls.

Basel / ECB Models

Capital-model lineage: strong on validation, narrower in scope, without a consumer or GenAI dimension.

What stands out: On AI, GenAI and consumer protection, the RBI draft is now among the most explicit prudential frameworks globally, Indian REs compliant with it will exceed most cross-border regulatory expectations

Reading It Correctly: Seven Misreadings to Avoid

Principles-based and proportionate guidance is open to misinterpretation. Seven points require clarification before an institution designs its response to the draft.

COMMON MISREADING	WHAT THE GUIDANCE ACTUALLY REQUIRES
Proportionality means lighter governance for AI.	Proportionality scales rigour to risk. High-impact AI attracts more scrutiny, not less, and no model is left ungoverned.
Buying a vendor or foundational model transfers the risk.	Accountability is non-delegable. The RE must independently validate and monitor third-party models regardless of any certification.
Only credit and capital models are really in scope.	All models with material decision impact are in scope, including rule engines, ML, GenAI and end-user tools.
Annual validation and attestation are sufficient.	Monitoring must be continuous and lifecycle-embedded; drift and limit breaches must be caught between formal cycles.
Explainability is a soft, best-effort expectation.	It is testable. Where explainability is limited, compensating controls and enhanced validation are required.
Human oversight can be a dashboard that can be reviewed later.	Oversight must be able to intervene —override, suspend and kill-switch with safeguards against automation bias.
Retired models can be deleted from the inventory	Decommissioned models must stay in the inventory for a minimum of 10 years, with lineage intact.

From Requirement to Operating Model: How Solytics Partners Helps?

Meeting these obligations at enterprise scale is an operating-model requirement, not solely a policy one. Solytics addresses it with two engines on a single control plane, governing traditional, ML, GenAI and agentic models under one accountable framework.

ENGINE 01



The intelligence and assurance engine, builds, validates, documents and assures every model, from statistical and ML through GenAI and agents.

BUILD • VALIDATE • DOCUMENT • ASSURE

ENGINE 02



The enterprise governance spine, the single, board-accountable register where every model is identified, tiered, governed and made audit-ready.

INVENTORY • TIER • GOVERN • EVIDENCE

Mapped to the Requirements: Governance Pillars

Governance Foundation

- Board-Approved MRMF
- Three Lines of Defence
- Enterprise Inventory
- Risk Tiering

Validation and Monitoring

- Independent Validation
- Approvals & Exceptions
- Monitoring & Drift
- Third-Party Models

AI / GenAI / Agentic AI

- Explainability & Fairness
- Hallucination Controls
- Human Oversight & Kill-Switch
- Disclosure & Audit

End-to-end eMRM Lifecycle Coverage

Identify > Develop > Validate > Approve > Deploy > Monitor > Govern > Assure

One Governed Estate

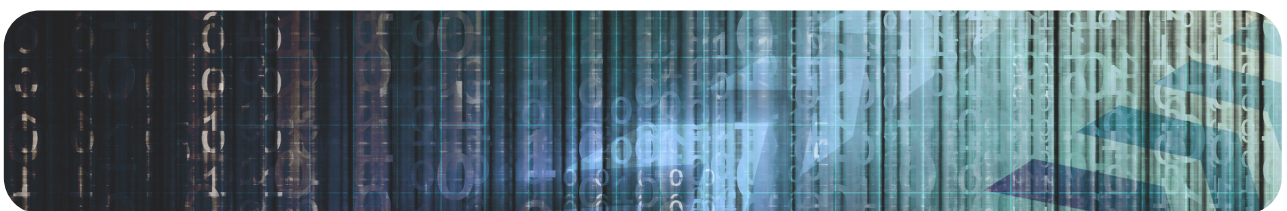
Every model, traditional, ML, GenAI and agentic, identified, tiered and controlled in a single board-accountable system.

Effort Aligned to Risk

Validation, approval and monitoring intensity scale with each model's tier, concentrating scrutiny where it matters.

Defensible on Demand

Audit-ready evidence and full lineage at every lifecycle stage, ready for supervisory and internal-audit review.



About Solytics Partners

SolyticsPartners is a global leader in AI governance, model risk management, and financial crime analytics, a Category Leader in the 2026 ChartisRiskTechQuadrant® for AI Governance, ranked consecutively for 3 years (2024-26) in the annual ChartisRiskTech100® rankings, recognized on PathFin.AI, MAS's regulator-curated AI marketplace, and winner of at the Regulation Asia Awards 2025.

Combining practitioner-depth domain expertise with advanced analytics and cloud-native engineering, Solytics serves leading financial institutions across the US, UK, MENA, APAC and Canada, helping them govern every model, from statistical and ML through GenAI and autonomous agents, in one regulator-ready framework.

The Solytics Partners AI Governance Ecosystem

MRM Vault - the enterprise governance spine: one board-accountable register where every model is inventoried, tiered, governed and made audit-ready.

Nimbus Uno - the intelligence and assurance engine that builds, validates, documents and assures every model from statistical and ML through GenAI and autonomous agents.

DISCLAIMER

This document is provided for informational purposes only and does not constitute legal, regulatory or professional advice. It reflects the authors' interpretation of the RBI's draft "Guidance on Regulatory Principles for Model Risk Management, 2026" (PR 2026-2027/528, released 24 June 2026), which is a draft issued for public consultation with comments due 24 July 2026. Provisions may change on finalisation, and readers should refer to the primary RBI text and seek independent counsel on application to their circumstances. SolyticsPartners makes no guarantee as to accuracy or completeness and disclaims any liability arising from use of this information. Product capabilities described are indicative and may vary by deployment and scope of work. All third-party framework names are the property of their respective owners.