



Artificial Intelligence Supervisory Framework

Nonbank AI Supplements

Conference of State Bank Supervisors

Version 1.0

Approval Date: August 13, 2026

Using the Nonbank AI Supplements 2

When to Use Each Supplement 2

Supplement Routing Guide..... 2

Using More Than One Supplement..... 3

Third-Party and Vendor Oversight..... 4

Model Risk Review 7

Consumer Protection 10



Using the Nonbank AI Supplements

The Nonbank AI Supplements are intended to help examiners apply AI-specific thinking when existing supervisory resources are used in areas such as third-party and vendor oversight, model risk review, and consumer protection review. They are not standalone examination procedures and do not replace existing supervisory resources, examiner judgement, or applicable legal requirements.

Examiners should begin with the Core Examiner Guide and use the Examiner Work Program to understand the purpose, source support, and examiner focus for each scoping question and procedure. When that review indicates a relevant existing supervisory area, examiners may use the applicable supplement as an overlay to identify AI-specific issues that may warrant additional review.

The supplements are narrow by design. They are intended to highlight issues that may not be fully captured by existing review resources, such as vendor-provided AI tools, embedded AI functionality, model or output limitations, changing performance over time, generative AI risks, sensitive data use, consumer-facing impacts, and reliance on AI-enabled recommendations or decisions.

Where AI use raises information security, cybersecurity, or technology control concerns, examiners may consider whether those issues should be reviewed through applicable IT, cybersecurity, third-party, or operational risk resources.

When to Use Each Supplement

Use the Nonbank AI Supplements after initial scoping or inventory review indicates that an institution's AI use may implicate an existing supervisory review area. The supplements may be used, as applicable, to support review within existing third-party and vendor oversight, model risk, or consumer protection resources.

The supplements are not mutually exclusive. A single AI use case may implicate more than one supplement. For example, a vendor-provided AI underwriting tool may raise third-party oversight, model risk, and consumer protection considerations. Examiners should use judgment to determine which supplement(s) is relevant based on the institution's AI use, risk profile, business model, consumer impact, and supervisory context.

Supplement Routing Guide

Supplement	Use when the review involves...	Relevant framework connections	Primary review focus
Third-Party and Vendor Oversight	Third parties, vendors, service providers, external platforms, AI-as-a-service, vendor-provided models, embedded AI functionality, or generative AI tools supported by third parties.	IS-4, IS-5, IS-8, GO-7, AI-6, GE-3, GE-4	Whether AI capabilities provided, hosted, supported, or embedded by third parties are identified, understood, and addressed within existing vendor oversight processes.
Model Risk Review	AI-enabled models, predictive tools, machine learning systems, scoring tools, decision-support tools, vendor-provided models, or AI outputs that materially support analysis, operations,	IS-3, IS-7, AI-2, AI-3, AI-4, AI-5, GO-6, GE-5, GE-6, GE-8	Whether model risk concepts such as intended use, limitations, validation, performance monitoring, effective challenge, explainability, output oversight,



	prioritization, fraud detection, or decisions.		and drift are relevant to the institution's AI use case.
Consumer Protection	Customer-facing AI, decision-support tools, eligibility or underwriting processes, pricing, servicing, collections, fraud review, complaints, marketing, communications, adverse action, generative AI outputs, or other activities that may affect consumers or consumer outcomes.	IS-3, IS-7, IS-8, AI-3, AI-4, AI-5, GE-5, GE-6, GE-8	Whether AI use may affect consumer treatment, outcomes, explanations, communications, sensitive data handling, or compliance with applicable consumer protection requirements.

Using More Than One Supplement

Examiners may use multiple supplements where the AI use case involves overlapping supervisory issues.

Examples may include:

- A vendor-provided AI tool used to support underwriting, or eligibility decisions may implicate **Third-Party and Vendor Oversight, Model Risk Review, and Consumer Protection**.
- A generative AI tool used by staff to draft consumer communications may implicate **Consumer Protection** and, if externally provided or hosted, **Third-Party and Vendor Oversight**.
- A machine learning tool used for fraud detection or transaction monitoring may implicate **Model Risk Review** and, if outputs affect account access, servicing, fees, or other consumer outcomes, **Consumer Protection**.
- An AI tool that processes sensitive customer or consumer information through a third-party platform may implicate **Third-Party and Vendor Oversight** and **Consumer Protection**.
- An AI system that can take actions with limited human direction may implicate **Model Risk Review, Consumer Protection, or Third-Party and Vendor Oversight** depending on the action, system access, consumer impact, and vendor involvement.



Third-Party and Vendor Oversight

Purpose

This supplement is intended to help examiners apply AI-specific thinking when using existing third-party and vendor oversight resources. It is meant to highlight AI-related issues that may not be fully captured by the existing third-party and vendor review resources, especially where AI capabilities are provided, supported, hosted, or embedded by third parties, vendors, service providers, or external platforms. This supplement is not a replacement for existing third-party review resources. It is a short supplement to help examiners ask better AI-specific questions in that review.

Examiner Use

Use this supplement when the institution relies on a third party, vendor, service provider, or external platform for AI capabilities, including embedded AI features, vendor-provided models, externally hosted tools, AI-as-a-service, or generative AI tools supported by third parties. Use it alongside existing third-party and vendor oversight resources to focus attention on AI-specific issues that may warrant added review, such as embedded AI functionality, AI-specific due diligence, vendor dependencies, model or tool changes over time, output oversight, and contingency planning.

Examiner Considerations

- **AI Identification:** Has the institution identified which vendor relationships involve AI capabilities, including embedded AI features, and does it understand how those capabilities are delivered and controlled?
- **AI Due Diligence:** Has the institution incorporated AI-specific considerations into due diligence, such as data provenance, data use practices, privacy, security, intellectual property, explainability limits, model changes over time, and vendor dependencies?
- **Roles and Responsibilities:** Does the institution understand the vendor's role, the institution's role, and any shared-responsibility boundaries for oversight, monitoring, incident response, and remediation?
- **Contract Terms:** Do contracts or service expectations address AI-specific issues such as data rights, usage restrictions, service changes, model updates, disclosure of material incidents, access to testing or audit information, termination, fallback planning, or options to pause, revert, or otherwise address identified issues?
- **Data Use:** Does the institution understand whether customer, consumer, confidential, sensitive, or institutional data may be used, retained, shared, or used to train, tune, or improve AI systems?
- **Vendor Changes:** Where vendor-provided AI systems are changed, updated, or replaced, does the institution understand how it is notified of material changes and whether continued reliability, suitability, or oversight expectations are assessed for the institution's use case?
- **Ongoing Monitoring:** Is the institution monitoring vendor AI performance, material changes, incidents, and control environment issues on an ongoing basis rather than relying only on upfront due diligence?
- **Output Oversight:** Where the vendor provides model outputs or AI-enabled recommendations, does the institution have a process to validate, challenge, benchmark, or otherwise oversee those outputs for its own use case?
- **Concentration and Contingency:** Has the institution considered concentration, over-reliance, or contingency issues if a critical AI vendor changes terms, degrades in performance, experiences an incident, or is no longer available?



- **Generative AI Vendor Risks:** If generative AI is involved, has the institution addressed vendor risks related to embedded GAI technologies, data exposure, content provenance, value-chain opacity, and third-party access to institutional content or prompts?

Key Sources

Source document	Specific section / concept	Key source language
SR 23-04 Third-Party Relationship Risk Management	Overview; responsibility for third-party activities; inventory; critical activities	The institution's use of third parties does not diminish its responsibility to operate in a safe and sound manner and in compliance with applicable laws and regulations. Maintaining a complete inventory of third-party relationships and periodically conducting risk assessments supports sound oversight. More comprehensive and rigorous oversight is appropriate for third party relationships that support critical activities or activities with the greater potential impact.
Treasury AI Lexicon	AI governance; AI system	AI governance is the set of organizational policies, rules, frameworks, roles, and oversight processes that direct how AI is adopted, developed, deployed, and monitored across the AI lifecycle.
CRI FS AI RMF Guidebook	GV-6 / GV-6.1 third-party AI risk management	Policies and procedures should address AI risks and benefits arising from third-party software and data and other supply chain issues. The organization should establish processes for evaluating and selecting third-party AI technologies, conducting AI-specific due diligence, documenting AI-related vendor risks, monitoring compliance, and managing concentration and fourth-party risk.
NIST AI 600-1	GOVERN 6.1 and 6.2 third-party GAI risk actions	Suggested actions include implementing a use-case-based supplier risk assessment framework; updating due diligence and procurement vendor assessments to include embedded GAI technologies, data privacy, security, and other risks; inventorying all third-party entities with access to organizational content; including contractual clauses that allow evaluation of third-party GAI processes and standards; establishing incident response plans for third-party GAI technologies; and establishing continuous monitoring of third-party GAI systems in deployment.
SR 26-2 / Revised Model Risk Guidance	Vendor products; intended use; limitations; testing; monitoring; contingency planning	Vendor products may pose unique challenges because expertise and key model components may be external or proprietary. Institutions should understand intended use, limitations, assumptions, testing results, customization choices, ongoing monitoring, outcomes analysis, and contingency options for vendor products, as applicable to the institution's use and risk.



Examiner Focus

Focus on whether the institution treats vendor-provided AI as a managed dependency rather than a black box. The key questions are whether the institution knows where third-party AI is being used, whether due diligence, contracts, and monitoring address AI-specific considerations, whether vendor outputs and changes can be meaningfully overseen, and whether fallback or contingency planning is considered for more significant uses. If AI is embedded in vendor tools but not identified, if material vendor changes are not communicated or assessed, or if the institution relies primarily on vendor representations, examiners may need additional follow-up to understand whether vendor-provided AI is subject to appropriate oversight for the institution's use case.



Model Risk Review

Purpose

This supplement is intended to help examiners apply AI-specific thinking when using existing model risk review resources. It highlights AI-related issues that may not be fully captured by ordinary model risk review, especially where artificial intelligence systems, predictive models, machine learning tools, vendor-provided models, or AI-enabled decision tools influence analysis, operations, consumer outcomes, or decisions. This supplement is not a replacement for existing model risk review resources. It is a short supplement to help examiners identify AI-specific concerns within that review.

Examiner Use

Use this supplement when the institution uses models or AI-enabled tools in ways that materially support analysis, operational decisions, customer outcomes, or business processes, or when vendor-provided or externally developed AI tools are used in significant activities. Use it alongside existing model risk review resources to focus attention on AI-specific issues that may warrant added review, such as intended use, validation, explainability limits, performance drift, human challenge, vendor opacity, and ongoing monitoring.

Existing model risk resources may be relevant where AI is used in models or model-like activities, including scoring, forecasting, classification, prioritization, underwriting, fraud detection, or other decision-support uses. However, some AI uses, including generative AI, large language models, agentic tools, or similar capabilities, may present risks that are not fully addressed through traditional model risk review. Examiners may use this supplement, together with the Core Examiner Guide and Examiner Work Program, to identify AI-specific considerations and determine whether additional review is appropriate based on the institution's use case and supervisory context.

Examiner Considerations

- **Model Scope:** Does the institution understand whether the AI tool or model falls within its model risk framework, or whether model risk concepts may be relevant even if the tool is not formally categorized as a model?
- **Intended Use:** Has the institution defined, as applicable, the intended use, business purpose, and decision context for the AI system, model, or tool?
- **Assumptions and Limitations:** Does the institution understand the model's assumptions, limitations, training data dependencies, performance boundaries, and the conditions under which outputs may become unreliable?
- **Performance Change:** Where machine learning or other AI techniques are used, has the institution addressed risks related to model drift, performance degradation, data changes, or changing real-world conditions?
- **Validation and Testing:** Where validation, testing, benchmarking, or other review is performed, is it tailored to the AI tool's actual use and risk, including where vendor-provided tools are used in significant activities?
- **Human Challenge:** Does the institution apply meaningful human review, challenge, escalation, or override processes where AI outputs inform decisions, prioritization, analysis, or other consequential actions?
- **Explainability:** Where explainability is limited, has the institution identified how that limitation affects use, oversight, consumer impact, or reliance on the outputs?
- **Ongoing Monitoring:** Does ongoing monitoring address AI-specific issues such as shifting inputs, emergent weaknesses, unintended outcomes, feedback loops, or degradation over time?



- **Vendor Models:** If the institution uses vendor-provided AI or model outputs, does it independently assess performance and limitations rather than relying primarily on vendor representations?
- **Generative AI Outputs:** Where generative AI is used in analysis or decision support, has the institution considered risks such as confabulation, overreliance, non-repeatable outputs, and use outside intended purposes?

Key Sources

Source document	Specific section / concept	Key source language
SR 26-2 / Revised Model Risk Guidance	Risk-based approach; model materiality; validation; vendor products; ongoing monitoring	Model risk management should be risk-based and tailored to an institution's model risk profile, size, and complexity. Models of greater materiality warrant more comprehensive oversight. Vendor products present unique challenges because expertise and key model components may be external or proprietary. The guidance emphasizes understanding intended use, assumptions, limitations, testing, customization choices, contingency options, and ongoing monitoring and outcomes analysis.
Treasury AI Lexicon	AI model; AI drift/decay; AI risk assessment	An AI model is a component of an information system that uses computational, statistical, or machine-learning techniques to produce outputs from inputs. AI drift or decay refers to performance degradation over time in real-world settings. AI risk assessment is a process for identifying, estimating, and prioritizing risks arising from the operation and use of an AI system.
NIST AI RMF 1.0	Trustworthiness; lifecycle; context and intended use; post-deployment monitoring	AI risk management should be applied in AI-system-specific contexts and throughout the AI lifecycle. Intended purposes, beneficial uses, deployment context, and business use should be understood and documented. Trustworthiness depends in part on validity and reliability, transparency, explainability, and accountability. Risk management should continue through deployment, operation, and monitoring.
NIST AI 600-1	Confabulation; human-AI configuration; information integrity; value chain and component integration	Generative AI can produce confidently stated but false content, create risks of overreliance, and involve non-transparent third-party components or integrations that affect downstream accountability and reliability. Risks and controls vary depending on model, system, and use-case context.

Examiner Focus

Focus on whether the institution treats AI-enabled models and tools as subject to meaningful performance review, validation, challenge, and monitoring when those tools influence analysis, operations, consumer outcomes, or decisions. The key questions are whether the institution understands intended use and limitations, whether review is scaled to materiality and risk, whether human oversight is meaningful, and whether vendor opacity or AI-specific



behavior may limit the effectiveness of ordinary model risk review. If AI outputs are relied on without challenge, validation is generic or limited, drift or changing inputs are not monitored, or the institution relies primarily on vendor representations, examiners may need additional follow-up to understand how model risk is being identified and managed for the institution's use case.



Consumer Protection

Purpose

This supplement is intended to help examiners apply AI-specific thinking when using existing consumer protection review resources. It highlights AI-related issues that may not be fully captured by ordinary consumer protection review, especially where artificial intelligence is used in customer-facing activities, decision support, communications, eligibility determinations, pricing, servicing, fraud review, or other activities that may affect consumers or consumer outcomes. This supplement is not a replacement for existing consumer protection review resources. It is a short supplement to help examiners identify AI-specific concerns within that review.

Examiner Use

Use this supplement, as applicable, when the institution uses AI in ways that may affect consumers directly or indirectly, including customer-facing tools, decision-support tools, eligibility or underwriting processes, servicing or collections activity, fraud tools, complaint handling, communications, marketing, or other activities that shape consumer treatment or outcomes. Use it alongside existing consumer protection review resources to focus attention on AI-specific issues that may warrant added review, such as opacity, explainability, notices, disclosures, adverse outcomes, fair lending, UDAP/UDAAP, proxy effects, data quality, automation bias, and overreliance on model or tool outputs.

Examiner Considerations

- **Consumer Impact:** Is AI used in customer-facing activities or in activities that support or influence decisions affecting consumers?
- **Decision Influence:** Does the institution understand how AI outputs may affect consumer treatment, outcomes, communications, or decisions, even where the tool is described as support only?
- **Specific Reasons:** Where AI informs adverse actions, denials, pricing, servicing decisions, or other consequential actions, can the institution identify and communicate specific and accurate reasons for those outcomes where required?
- **Transparency:** Has the institution considered whether AI use may reduce transparency, obscure reasoning, or weaken explainability for employees, reviewers, or consumers?
- **Data and Proxy Risks:** Has the institution considered whether data used by the AI system, including proxies or alternative data, may contribute to inaccurate, inconsistent, or potentially unfair consumer outcomes?
- **Outcome Monitoring:** Does the institution monitor for differences in performance or outcomes across products, channels, customer groups, or use cases that may indicate consumer protection concerns?
- **Consumer Communications:** Where AI is used in customer communications, does the institution review generated or AI-assisted content for accuracy, appropriateness, and consistency with law, policy, and consumer-facing requirements?
- **Generative AI Risks:** If generative AI is used in consumer interactions or decision support, has the institution considered risks such as confabulation, misleading outputs, overreliance by staff, and the use of vague or inaccurate explanations?
- **Human Review:** Does the institution apply meaningful human review, escalation, or override processes where AI outputs influence consumer-facing actions or decisions?



- **Vendor Consumer Impacts:** Where vendor-provided AI is used in consumer-related activities, does the institution independently understand and oversee how that AI may affect consumer outcomes rather than relying primarily on vendor representations?

Key Sources

Source document	Specific section / concept	Key source language
ECOA / Regulation B	Adverse action notices; statement of reasons	Where adverse action notice requirements apply, creditors must provide a statement of specific reasons for the action taken. The statement of reasons must be specific and indicate the principal reason or reasons for the adverse action.
NAIC Model Bulletin	Section 3 expectations; unfair discrimination, inaccuracy, transparency, and AI-specific controls	The bulletin emphasizes that consumer-impacting decisions or actions made or supported by AI should be reviewed in light of applicable laws and regulations. It identifies risks such as inaccurate, arbitrary, capricious, or unfairly discriminatory outcomes and highlights the role of AI-specific controls in mitigating adverse consumer outcomes.
NIST AI RMF 1.0	Trustworthiness characteristics; fair with harmful bias managed; explainable and interpretable; accountable and transparent	Trustworthy AI includes characteristics such as valid and reliable, accountable and transparent, explainable and interpretable, privacy-enhanced, and fair with harmful bias managed. AI risks can emerge from technical aspects combined with societal factors related to how a system is used and the context in which it is deployed.
NIST AI 600-1	Confabulation; harmful bias; human-AI configuration; information integrity	Generative AI can produce confidently stated but false content, amplify harmful bias, increase risks of overreliance or automation bias, and lower the barrier to generating misleading content. These risks can vary based on the model, system, and use-case context.

Examiner Focus

Focus on whether AI use may affect consumer treatment, outcomes, explanations, or communications in ways that are not fully visible through ordinary process review. The key questions are whether the institution understands where AI influences consumer-related activity, whether specific and accurate reasons can be identified and communicated where required, whether accuracy, explainability, fair lending, UDAP/UDAAP, or other applicable consumer protection concerns are considered, and whether human oversight remains meaningful where AI is used in contexts with greater potential consumer impact. If AI-supported decisions or communications cannot be explained clearly, if vague or inaccurate reasoning is used in place of actual factors, if outputs are relied on without challenge, or if performance, proxy effects, outcome differences, or generative AI errors have not been considered, examiners may need additional follow-up to understand how consumer protection risks are being identified and managed for the institution's use case.