

Generative & Agentic AI — Risk & Compliance Review Question Guide

Illustrative example — not official guidance or advice. This is a generic framework built to show how AI-specific review questions can supplement a traditional product-risk review. It is not exhaustive and does not represent the policy of any institution; adapt it to your own organization’s frameworks and risk appetite.

Version history

Version	Date	Summary of changes
v1.0	2026-08-06	Initial draft — de-identified public version.
v1.1	2026-08-06	Reframed original criteria as questions; retitled column to “Product Review Questions”; added ISO/IEC 42001 references.

About this guide

This guide is a companion to a risk and compliance review process. For each specialist review group, it sets out the questions used to determine whether that group’s review is required for a new product or product change. In addition, for each group it provides a supplemental set of review questions specific to **Generative AI and Agentic AI** use cases — the risk considerations a conventional product review is not designed to surface, such as model behavior and data, decision explainability, the degree of autonomy an agent has, human oversight, and accountability. These AI questions are additive to — not a replacement for — a group’s original criteria: apply them whenever a change is triggered under a group’s criteria and involves GenAI or Agentic AI.

How to use this guide

- Use it alongside your existing review guidelines — the AI questions are additive, not a replacement. If a group is triggered under its original criteria and the change uses GenAI/Agentic AI, ask the AI questions too.
- “Generative AI” = systems that produce content (text, code, images, decisions) from foundation/large language models. “Agentic AI” = AI that plans and takes actions autonomously across tools and systems (it can call APIs, move data, and in some designs move money or change accounts).
- The Foundational questions below apply to every group. The table then adds group-specific questions.

Foundational (cross-cutting) AI questions — ask for every review

- Model & provenance: What model powers this (in-house, fine-tuned, or third-party foundation model/API)? Who hosts it and where does prompt/response data flow?
- Data: What data trains, fine-tunes, or grounds (RAG) the model? Does it include non-public personal information (NPI), confidential, or biometric data? Is our data excluded from vendor training and logging?
- Human oversight & autonomy: Is there a human in/over the loop? For agentic use, what actions can the agent take autonomously, and where are the approval gates, limits, and a kill switch?
- Accuracy & explainability: How are hallucination/confabulation, accuracy, and drift measured? Can an output or decision be explained to a customer, auditor, and regulator?
- Security & abuse: Is the system tested against prompt injection, jailbreaking, and data exfiltration? Are agent tool/system permissions scoped to least privilege?
- Accountability & audit: Are prompts, outputs, and agent actions logged in a tamper-evident, reviewable trail? Who is the accountable owner if the model is wrong?
- Governance mapping: Has the use case been assessed against the bank’s model risk (SR 11-7) and AI governance standards (e.g., NIST AI RMF and its Generative AI Profile, and ISO/IEC 42001)?

References used to frame the AI questions: NIST AI Risk Management Framework and its Generative AI Profile (NIST AI 600-1); emerging NIST/CSA agentic-AI profile work; Federal Reserve/OCC SR 11-7 model risk management; CFPB guidance requiring specific adverse-action reasons for AI/ML credit models under ECOA/Reg B; and current agentic-AI-in-financial-services governance literature; and ISO/IEC 42001, the AI management system (AIMS) standard.

Risk and Compliance Specialist Review Table

Review Group / Risk Domain	Product Review Questions	Generative & Agentic AI Review Questions
Accessibility	● Will the change result in a new customer-facing website or mobile application? ● Does the change impact how a customer will use the website or mobile application — features and functionality? ● Will the change introduce a new digital platform (e.g., ATM touchscreen, website, mobile app)? ● Will there be a call center or telephone interaction with customers?	● If the change adds a customer-facing GenAI feature (chatbot, virtual/voice assistant, generative UI), has dynamically generated content been tested for accessibility (WCAG 2.1/2.2) and screen-reader compatibility? ● Is there an accessible, non-AI alternative or human hand-off for customers who cannot use the AI interface (speech, hearing, cognitive, or motor disabilities)? ● For conversational/voice agents: are pacing, timeouts, and re-prompts adjustable, and are text alternatives provided for audio-only interactions? ● Can a customer using assistive technology complete the same agentic task end-to-end (not just start it)? ● Do AI-generated images, charts, or documents include accessible alt text / structure?
Accounting Policy	● If there is Revenue Recognition: ○ Type of arrangement (loan, lease, guaranty, revenue share, other); single vs. multiple agreements; same party; new agreement or amendment. ○ All parties; any third parties; single vs. multiple deliverables; gross vs. net revenue/expense; timing of recognition; costs — can they be capitalized? ● Loan & loan-origination accounting and classification: ○ Costs/fees; reserves; reporting. ● Troubled Debt Restructuring — could the change result in TDR classification, and can it be flagged/accounted/reported correctly? ● For Investments is there — accounting, classification, reporting? ● Are there Guarantees, commitments, other liabilities/contingencies? ● Variable Interest Entities & consolidation (variable interest held? primary beneficiary? new/revised contracts). ● Lease accounting — type of lease and appropriate treatment?	● Are the costs to develop, train, fine-tune, or license the GenAI/agentic solution (compute, foundation-model/API token fees, data labeling, integration) properly classified as capitalizable software development vs. period expense? ● Does a third-party model arrangement (usage/token-based pricing, revenue share, committed capacity) create new revenue-recognition, expense-timing, or gross-vs-net presentation questions? ● Could autonomous agent actions (e.g., dynamic pricing, automated fee assessment/waivers) change the timing or measurement of revenue recognition? ● Does the solution create a new intangible asset (proprietary model, curated dataset) requiring recognition/impairment consideration? ● If the agent influences loan pricing, fees, or modifications, could it drive a classification (e.g., TDR) that must be flagged and reported correctly?

Review Group / Risk Domain	Product Review Questions	Generative & Agentic AI Review Questions
	● Derivative considerations; ○ Are Digital assets/currencies involved? ○ Are alternative rate indices (LIBOR transition) impact?	
Controller's Operations / Reconciliation	● General ledger impacts — new interface, legal entity, or other impacts requiring GL setup/modification? ● Reconciliation impacts? ○ Which areas reconcile (or don't)? New accounts for reconciliation? New settlement processes? Regulatory-capital (Basel) impacts? Technical/file-transmission impacts? ● Fixed assets impacts? — most follow existing BAU booking/depreciation. ○ Any asset considerations differing from policy? Impact to assets already on the books (modification, write-off)? ● Payables Impact?: ○ Payment needs differing from policy? Check production or modifications facilitated on behalf of a business line?	● Will the GenAI/agentic solution create new GL interfaces, automated journal entries, or settlement processes that require reconciliation controls? ● If an agent posts or initiates transactions autonomously, how are those entries reconciled, and is there a complete, tamper-evident audit trail of each agent action and the data it relied on? ● Are new suspense/clearing accounts needed to capture AI-driven exceptions, reversals, or errors? ● Do AI-driven payment or check-production requests differ from policy, and are they subject to human authorization limits? ● How are AI output errors detected and corrected before they flow to the ledger?
Tax	● Are there Information-reporting requirements (e.g., 1099/1098/1042)? ● Are there withholding tax requirements? ● Are there changes to the general ledger — new accounts, account usage, or source systems. ● Are there new legal entities, jurisdictions, or changes to how the business goes to market? ● Are there potential sales/use tax, VAT, GST, gross receipts or transactional taxes on new purchases/acquisitions/divestitures? ● Are there new/changed cross-border transactions (services, equipment, software licenses, multijurisdictional contracts) that may impact transfer pricing? ● Are there changes to tax-sensitive products/business lines (e.g., equipment finance, leasing, mortgage)? ● Are Digital assets/currencies? ● Are alternative rate indices (LIBOR transition) impacted?.	● Does use of a third-party/cloud GenAI service create new sales/use tax, VAT/GST, or transfer-pricing considerations (e.g., cross-border data processing, intercompany charges for a shared AI platform)? ● Could AI-generated documents or determinations affect information-reporting or withholding accuracy — and are those outputs controlled/validated? ● Does deploying or hosting the model in a new jurisdiction create nexus or a new taxable presence? ● Are new GL accounts or source-system changes introduced by the AI platform that Tax must map? ● If an agent takes tax-relevant actions (invoicing, fee assessment) autonomously, are the tax outcomes reviewable and correctable?
Accounting & External Reporting	● Chart of accounts — impact to the balance sheet/income statement and application changes identified? ● Legal entity — clear identification of the entity offering the new product/change?	● If AI outputs feed regulatory or financial reports, what controls ensure the accuracy, completeness, and traceability of AI-generated or AI-classified figures? ● Are the financial-statement effects of the AI initiative understood and disclosed — development costs (capitalized vs. expensed), efficiency-

Review Group / Risk Domain	Product Review Questions	Generative & Agentic AI Review Questions
	● Intercompany — review for intercompany dependencies? ● Reconciliation — sound practices in place; account certifier identified? ● Accounting flow of transactions — GL entries and flow understood? ● GL changes (new midlines / profit-plan lines)? ● Material change in business volume requiring new explanations (SEC filings, variance analysis)? ● Financial-statement impact (net interest margin, efficiency ratio, etc.)? ● Regulatory-capital (Basel) impact — RWA and capital; ensure credit exposures can be risk weighted? ● Impacts to regulatory filings; product-code changes to ensure correct classification?	ratio impact, and any material AI model-risk disclosures? ● Could AI-driven classification or product-coding introduce misstatement risk (e.g., in RWA or product classification) that must be independently checked? ● Is there a documented owner and reconciliation/certification process for any account or report populated by AI output? ● Does a material change in AI-driven business volume require new management explanations or variance analysis?
Business Continuity Planning	● Does the change directly, or indirectly, impact a mission-critical process?	● Is the GenAI/agentic system — or its third-party model provider/API — part of, or a dependency of, a mission-critical process? ● What is the fallback if the model, API, or agent is unavailable, degraded, or produces unsafe output — is there a manual or reduced-function mode? ● Is there a tested kill switch to halt an autonomous agent, and a defined process to reverse or contain actions it already took? ● What are the RTO/RPO and single-vendor concentration risks for the foundation-model provider (and its sub-providers)? ● Are model/version updates by the vendor treated as changes that could disrupt the process?
Financial Crimes (AML / Sanctions)	● Does the change target higher-risk customer types (e.g., non-bank financial institutions, NGOs, professional service providers, government-related parties)? ● New or changed money-movement flows, transaction types, processes, payment systems, or movement to foreign jurisdictions? ● Alters/creates data feeds, customer-info systems, or tech related to CIP/KYC, risk scoring, sanctions screening, or transaction monitoring (incl. volume increases)? ● Jurisdictional elements foreign to where the business is offered/booked? ● Will CIP/KYC, sanctions screening, or transaction monitoring be performed by an external party on the institution's behalf? ● Cash-intensive businesses or large cash amounts triggering	● If GenAI/agentic tooling performs or supports CIP/KYC, sanctions screening, risk scoring, or transaction monitoring, has the model been validated for false-negative/false-positive rates, and can each alert or decision be explained and evidenced to regulators? ● Could an autonomous agent initiate or facilitate money movement — and what controls prevent it from being manipulated (prompt injection, data poisoning) to bypass AML/sanctions controls? ● Is customer/transaction data sent to a third-party model, and does that create data-jurisdiction or sanctions exposure (where is it processed/stored)? ● Does the AI create new typologies or reduce the auditability/traceability of alerts and dispositions? ● Are AI-generated suspicious-activity narratives or investigative summaries reviewed by a qualified human before filing?

Review Group / Risk Domain	Product Review Questions	Generative & Agentic AI Review Questions
	regulatory reporting (e.g., currency-transaction reports, monetary instruments, foreign-account reporting)? ● Are Digital assets/currencies impacted?	
Compliance	● Compliance reviews all change requests to determine compliance impact based on regulations that may be impacted.	● Does the GenAI/agentic use case implicate unfair-or-deceptive-practices (UDAAP), fair lending, or disclosure rules — and is there a documented compliance assessment of the AI's outputs and decisions? ● Is there a mechanism to detect, review, and remediate non-compliant AI-generated content or decisions before or after customer impact? ● Are AI-generated customer communications/disclosures reviewed for accuracy and non-deception? ● Has the use case been mapped to applicable AI-specific rules (e.g., adverse-action expectations, state AI laws, EU AI Act if in scope)? ● Is there ongoing monitoring/testing of the model for compliance drift as it is updated?
Credit Risk	● Introduction of credit/transaction products and services; portfolio purchases; acquisitions? ● Modification to existing products/services (e.g payments, ACH, loan/lease, installment, mortgage, wealth)? ● Service providers with credit risk from delay of funds; changes to underwriting/account-decisioning criteria? ● Changes moving funds in/out of the institution (e.g., wires)? ● New/modified account-acquisition channels; ATM and debit-card changes? ● Model or tool related changes; default/collections changes? ● Pricing/fee changes affecting revenue reporting. ● Third-party lending relationships; changes to any credit-processing systems (approval/origination systems, data warehouses). ● Are Digital assets/currencies impacted?. ● Changes to terms/nature that alter underlying credit-risk characteristics?	● Does the GenAI/agentic tool influence underwriting, credit or account decisioning, pricing, line management, or collections? If so, is it inventoried and validated under model risk governance? ● Can each credit decision be explained with specific, accurate reasons — sufficient to generate compliant ECOA/Reg B adverse-action notices (regulators have signaled that generic reasons are not enough for complex AI/ML models)? ● Has the model been tested for disparate impact / fair-lending risk and proxy discrimination across protected classes? ● Could an autonomous agent extend, modify, restructure, or collect on credit without human approval — and what limits, dual-control, and approval gates apply? ● Is there ongoing drift and performance monitoring so credit-risk characteristics don't shift unmonitored as the model or data changes? ● For third-party/foundation models, what validation and transparency are possible without access to model internals?
Customer Experience	● Customer impact at scale — or a pilot involving customers? ● Change impacts how customers interact with us (digital navigation change, or must use a different channel)? ● Additional effort required by the customer to complete an action (extra steps/mediums)? ● Eliminating a previously offered product/service (replacement,	● Are customers clearly informed when they are interacting with AI rather than a human, and is there an easy, low-friction path to reach a human? ● Have AI-generated responses been tested for accuracy, tone, and the risk of hallucinated or misleading guidance that could harm or confuse the customer? ● Does the agentic flow add or remove steps for the customer, and has

Review Group / Risk Domain	Product Review Questions	Generative & Agentic AI Review Questions
	account closure)? ● Updating manual or digital processes that impact customers; launch of a new product/service? ● Change requires digital or paper communication/notification to customers, or responds to a regulatory change affecting customer interaction?	the failure experience (wrong answer, agent stuck, refusal) been explicitly designed? ● Is AI personalization free of manipulative “dark patterns,” and can the customer opt out of AI handling? ● How is customer sentiment/complaint signal monitored specifically for AI interactions?
Data Protection / Privacy	● How is data-handling, storage, and technology-risk implications of the change? ● Are Digital assets/currencies impacted?	● Is non-public personal information (NPI) sent to, processed by, or retained by the GenAI model or third-party provider — and is our data contractually excluded from vendor training, tuning, and logging? ● Are prompts and outputs that may contain customer data logged; if so, how are those logs protected, access-controlled, and retention-limited? ● Does the model create new inferences, profiles, or derived data about customers that constitute a new use of personal data requiring assessment? ● Is data minimization enforced (only the minimum necessary fields reach the model), and is sensitive data masked/tokenized before inference? ● Where is prompt/response data physically processed and stored (data residency), and does that meet privacy requirements?
Digital Assets Operations Risk	● Are Digital assets or digital currencies involved.	● If an agent can read from or transact against digital-asset systems, what constraints, approval gates, and limits govern autonomous execution? ● Are AI-driven digital-asset operations fully logged and reversible, with a stop capability? ● Is the model hardened against manipulation that could trigger erroneous digital-asset movements?
Data Governance	● Data sourcing (customer data): does the change use customer data outside approved/authorized sources? ● Data sourcing (authorized data): will it create any new authorized data source or system of record? ● Data quality: known data-quality issues with sources? Data-quality checks/processes being built in? ● Data usage: any ethical concerns? Any new data types, attributes, domains, or data products created? ● Digital assets/currencies.	● What specific data sources train, fine-tune, or ground (RAG) the model — and are they all authorized? Does the model rely on data outside approved sources? ● Are input data-quality controls in place (to prevent “garbage-in”), and is output quality/accuracy/drift monitored over time? ● Does the model create new data types, embeddings, inferences, or “data products,” and are these catalogued with lineage and governed for permitted use/consent? ● Are there ethical-use concerns with the training or inference data (secondary use, sensitive attributes, consent scope)? ● Is there end-to-end lineage/traceability from training and prompt data

Review Group / Risk Domain	Product Review Questions	Generative & Agentic AI Review Questions
		through to model outputs used in decisions?
Fraud Risk	- Does new/existing product changes, features, or services that may expose the institution to fraud (internal/external), including: - New or changed criteria for access to accounts/portals/websites (view/obtain/change info or execute a transaction): - Authentication - Authorization - New/changed criteria for moving money internally or externally. - Elder / vulnerable-adult financial exploitation; account take-over (ATO) risk. - ID theft, synthetic ID, first/third-person fraud; fraud education/content. - Biometric/authentication changes.	- Does the GenAI capability create new fraud vectors — deepfake/voice-clone social engineering, synthetic-identity generation, AI-assisted phishing, or prompt-injection to move money? - If an agent can move money or change account access/credentials, what step-up authentication, authorization limits, dual control, and anomaly detection constrain it? - Could generative outputs be weaponized against customers (convincing phishing, fake notices) — and is AI-generated customer content monitored/watermarked? - Is the model itself a fraud target (data poisoning, adversarial inputs, model theft), and how is that mitigated? - Does the change weaken existing authentication (e.g., voiceprint) against AI-based spoofing?
Enterprise Privacy Office	- Acquiring NPI from a third party or sharing NPI with a third party? - Transferring NPI outside the direct control of the institution's personnel or facilities? - Will biometric information be used for identification or authentication? - Will NPI be used to develop or train a model?	- Will NPI or biometric data be used to develop, train, fine-tune, or prompt the model (including via RAG or logs)? - Does sending prompts/data to a third-party foundation model constitute sharing NPI with, or transferring it outside the control of, the institution — and are data-processing/contract terms (no training on our data, deletion, sub-processor limits) in place? - Are privacy notices, consents, and permitted-use scopes updated to cover AI processing and any new inferences the model produces? - Do new AI inferences about individuals trigger privacy, profiling, or automated-decision obligations? - Is there a documented privacy impact assessment (DPIA) for the AI use case?
Fair & Responsible Banking	- Is there any change relating to/involving/impacting a credit, deposit, or payment-processing product or service — including pricing/fee changes and add-on/ancillary products? - Applies across processes: product/service development, marketing & sales, account opening, servicing, default management; and non-English external communications. - May include Deposits, Distribution, Commercial Lending, Corporate Services, Consumer and Business Lending/Leasing, or Issue and Acquiring - May include Product/ Service Development, Marketing and Sales, Account Opening, Servicing.	- Does the GenAI/agentic system influence marketing, targeting, pricing, offers, credit, deposit, or payment decisions in a way that could create disparate impact, redlining, or steering risk? - Is model output tested for bias across protected classes, and is AI-driven personalization/targeting checked for proxy discrimination (e.g., ZIP, device, behavioral proxies)? - For AI credit/eligibility decisions, can specific, accurate reasons be produced (fair-lending explainability) and are adverse-action notices compliant? - Are AI-generated marketing/disclosures accurate, non-deceptive, and

Review Group / Risk Domain	Product Review Questions	Generative & Agentic AI Review Questions
		non-discriminatory (UDAAP), including any non-English AI translations? ● Is there a documented fairness assessment and ongoing monitoring for the AI use case?
Human Resources	● Does the change involves any policy change or potential violation of state/federal HR legal or compliance requirements, including but not limited to: ○ Diversity & inclusion; accommodations; payroll, compensation, incentives & benefits; recruiting & onboarding. ○ Employee leave policies; employee records; workplace harassment; employee benefits; use of employee data; physical safety of employees.	● If GenAI is used for recruiting, screening, performance, compensation, or workforce decisions, does it comply with employment law and AI-in-hiring rules (e.g., EEOC guidance, NYC Local Law 144 bias audits, state AI laws), and has it been bias-audited? ● Does the change use employee data to train or fine-tune models, and is that a permitted, disclosed use? ● Could agentic automation materially change roles, incentives, or headcount in ways requiring HR/change-management and legal review? ● Is there human oversight of AI-influenced employment decisions and an appeal path for employees? ● Could the AI create workplace-safety or surveillance concerns for employees?
Information Security	● Will the change introduce new functionality in a new or existing application? ● Will the change alter the architecture and design of an information system? ● Will the change introduce new technology such as AI, cloud, API, or other technology? ● Will the change establish/alter electronic transfers of non-public data outside the institution's network? ● Will the change impact security standards/protocols for data storage or usage on/off premises? ● Will the change provide/alter standard or remote access rights to internal systems or the network? ● Does the change impact Digital assets/currencies?	● Where is the model hosted and where does prompt/response data flow (third-party API, cloud region), and is that data encrypted in transit/at rest and logically segregated? ● Is the system tested against AI-specific threats — prompt injection, jailbreaks, insecure output handling, training-data poisoning, model/data exfiltration, and sensitive-info disclosure (e.g., OWASP LLM Top 10)? ● For agentic systems, are the agent's tool/system permissions scoped to least privilege, with guardrails, allow-lists, and rate limits on the actions it can take? ● Are all AI inputs/outputs and agent actions logged and monitored for security events, and are API keys/model artifacts/secrets protected? ● Is there a secure-by-design review and a defined incident-response path for AI-specific incidents?
Legal	● Legal reviews all change requests to determine any legal impact.	● Who owns, and who is liable for, AI outputs and autonomous agent actions — and are third-party model terms, indemnities, warranties, and liability caps understood? ● Are there intellectual-property risks in training data or generated outputs (copyright, licensing, trade secret), and who owns prompts/outputs? ● Does the use case trigger AI-specific regulation (e.g., EU AI Act high-risk obligations if in scope, state AI/automated-decision laws, UDAAP)? ● Are contracts and disclosures updated for AI use, data rights, and

Review Group / Risk Domain	Product Review Questions	Generative & Agentic AI Review Questions
		consumer notice/consent? ● Is there a defensible record (documentation, testing, human oversight) to support the AI decision if challenged?
Market Risk	● Any new/modified product/service or change involving trading businesses, including: derivatives; foreign exchange; loan trading; fixed income; municipals; mortgage servicing rights; pipeline hedging; MSR valuation.	● If GenAI/agentic tools are used in pricing, trading, hedging, valuation, or research, are their outputs independently validated, and is there mandatory human review before any market-facing action? ● Could an autonomous agent place, modify, or cancel trades or hedges — and what position limits, pre-trade checks, circuit-breakers, and kill switch apply? ● Are AI valuation/forecast models governed under model risk, including testing for instability and regime change? ● Are there controls to prevent AI-driven herding or correlated errors across desks?
Model Risk	● Does the change use a quantitative method (model/tool) whose output is a predicted, forecasted, or estimated amount, rate, risk score, price, or value? ● Does the change use a quantitative method/approach applying statistical, economic, financial, or mathematical theories to process input into estimates/predictions? ● Does the change involve a software application/platform with an embedded quantitative method? ● Does the change uses machine learning or artificial intelligence?	● Is the GenAI/foundation model inventoried as a model and validated — including for non-deterministic output (benchmarking, hallucination/accuracy rates, robustness), not just a point estimate? ● How are generative outputs (text/decisions) tested, since traditional back-testing of a single number may not apply — what metrics and human evaluation are used? ● For agentic systems, is the agent policy/decision logic (tool use, planning, autonomy limits) governed as part of the model, not just the underlying LLM? ● For third-party/foundation models, what validation, documentation, and ongoing monitoring are feasible without internal access, and how are vendor model updates controlled? ● Is ongoing monitoring in place for drift, degradation, and emergent behavior, with defined thresholds and fallback?
Operational Risk	● Does the change create new processes/procedures or require updates to existing ones?	● Does the agentic system create new processes/procedures that need documented controls, defined human-in-the-loop checkpoints, and exception handling? ● What is the operational failure mode if the agent takes a wrong action autonomously — how is it detected, contained, reversed, and escalated? ● Is there a runbook to disable the agent and roll back its actions, and are staff trained on AI failure handling? ● How are AI errors, overrides, and near-misses captured as operational-risk events and fed back into controls? ● Are process controls resilient to AI variability (same input, different

Review Group / Risk Domain	Product Review Questions	Generative & Agentic AI Review Questions
		output) and to model/version changes?
Payments Risk	● Are there changes to the transfer of monetary value to/from the institution's owned or domiciled accounts, domestic or international? ● Are Payment systems impacted? This may include real-time payments, wire (e.g., SWIFT/FedWire), ACH, and instant-payment networks. ● Does the change establish a new, or involve/affect an existing, product/service/process/software related to payments or a payment system? ● Does the change involve seeking or ceasing participation in a payment system?	● Could a GenAI/agentic system initiate, approve, route, or modify payments/wires/ACH autonomously — and what authorization limits, dual control, velocity/anomaly checks, and payee-verification apply? ● Is the payment agent hardened against prompt injection or manipulation that could redirect funds or change beneficiaries? ● Are all AI-initiated or AI-influenced payments fully logged, reconcilable, and reversible where possible, with a real-time stop/kill capability? ● Does the AI change fraud or misdirected-payment exposure (e.g., business-email-compromise-style manipulation), and are limits set below escalation thresholds? ● Is there human confirmation for high-value or novel payment instructions generated by AI?
Physical Security	● Physical transport of confidential data outside the facility on a regular basis? Physical move of workspace/building? ● Space needing to be secured / additional equipment (cameras, alarms, card access)? Confidential data stored off-site? ● Affects physical handling/processing/storage of confidential info, cash, or negotiables? New risk to safety of employees/customers? ● Adverse event damaging reputation related to physical-security gaps?	● Does the AI change move confidential data off-premises to a third-party model/host in a way that changes physical data-handling/storage assumptions? (coordinate with Information Security / Privacy) ● If AI controls or informs physical-security systems (badge access, surveillance analytics), what human oversight and fail-safe exist against false accept/reject? ● Could an AI-driven physical-security failure create a safety or reputational event? ● (Often limited applicability — assign primarily where AI touches facilities, surveillance, or off-site data handling.)
Reputational Risk	● Are there new/expanded/changed incentives or incentive plans; impacts how employees perform their job? ● Are there environmental, social-responsibility, community, or equity impacts? ● Does the new product/service impact customers; external communications, media, social, or press (excl. customer comms)? ● Are there Co-branding, partnerships, JV, or strategic alliance; impacts a sensitive relationship category (e.g., cryptocurrency, project financing)? ● Does the change have customer impact? ○ at scale or a pilot; ○ new technology/system impacting customers; ○ expansion beyond current footprint; ○ eliminating a product/service.	● Could AI outputs (biased, offensive, hallucinated, or “a machine deciding about customers”) create reputational, ESG, or public-trust harm if they surface publicly? ● Is the use of AI disclosed appropriately, and does it align with the institution's public responsible/trustworthy-AI commitments and principles? ● Could an autonomous agent take a customer-visible action that, if wrong, becomes a media/social incident — and is there a rapid-response plan? ● Does the AI involve sensitive domains (crypto, surveillance, employment) that heighten reputational scrutiny? ● Is there monitoring for AI-related complaints, social sentiment, and press exposure?

Review Group / Risk Domain	Product Review Questions	Generative & Agentic AI Review Questions
Sales Practices	● Are there changes to risk/controls for sales practices by employees or third parties, especially account-opening controls and incentives/commissions? ● Does the change modify how the product/service is offered to consumers? ○ changes to how affirmative consent is captured; ○ changes to an existing product's marketing/selling/advising/account-opening; ○ changes to delivery via third-party vendors, e-media, or new channels/geographies.) ● Does the change modify the sales incentive/commission structure associated with the product/service?	● Does the GenAI/agentic system recommend, offer, cross-sell, advise on, or open products/accounts — creating risk of unsuitable recommendations, unauthorized account opening, or undue influence/dark patterns? ● How is affirmative consent captured and evidenced when an AI agent conducts the interaction, and is there human review before binding customer actions? ● Are AI-driven sales interactions monitored for suitability, disclosure, and compliance, and are incentives/commissions tied to AI-driven sales appropriately controlled? ● Could the AI steer customers toward higher-fee or higher-incentive products (conflicts of interest)? ● Is AI-generated sales/advice content reviewed for accuracy and non-deception?
Strategic Risk	● Will the change require investment-committee approval for capital expenditure? ● Is the change designated as a key/strategic business change per policy?	● Is the GenAI/agentic initiative aligned with the enterprise AI strategy and risk appetite, and — given emerging-tech and reputational exposure — should it be treated as a key business change requiring executive/investment-committee review? ● What is the strategic dependency/lock-in and concentration risk on a specific foundation-model provider, and is there an exit/substitution plan? ● Does the initiative create or erode durable competitive advantage, and are the scaling costs (compute/API) understood? ● Is there executive accountability and board-level visibility for material AI initiatives?
Technology Systems Risk	● Does the change impact system-change flags? ● Is a architecture-decision review required? ○ new technology/system; ○ technology upgrade/change; ○ complex/significant technology changes. ● Are there impacts to technology/platforms incl. custom software development (internal/external)? ● Does the change include: ○ AI, machine learning, ○ other emerging tech; ○ use of an API; ○ implementation of a cloud solution; ○ third-party vendors/hosting.	● Is the AI architecture — model hosting, orchestration, agent framework, tool integrations, and guardrails — documented via an architecture review, including non-deterministic behavior, latency, cost, and scalability? ● Are agent tool/system permissions scoped to least privilege with guardrails, allow-lists, and monitoring, and is there resilience/rollback for model or version changes? ● How are model/version updates (including silent vendor updates) tested and controlled before reaching production? ● Is there observability for AI (prompt/response logging, action traces, performance and drift dashboards)? ● Are third-party model/hosting dependencies inventoried, and is there a defined failover/degraded mode?

Review Group / Risk Domain	Product Review Questions	Generative & Agentic AI Review Questions
Treasury	● Does the change impact pricing of on/off-balance-sheet products? ● Movement of funds in/out/through the institution? ● Impact to the balance sheet or income statement? Regulatory-capital (Basel) reporting change / re-categorization of loan or deposit accounts? ● Change in product codes or general ledgers? ● Does the product supports foreign currencies? ● Material impact to deposit rate/fees that could drive large outflows?	● If GenAI/agentic tools inform pricing of balance-sheet products or the movement of funds, are outputs validated and is there human oversight before any balance-sheet or liquidity impact? ● Could an autonomous agent affect funding, deposit pricing, or liquidity movements — and what limits, approvals, and monitoring apply? ● Does AI-driven activity change deposit behavior/liquidity or off-balance-sheet commitments in ways Treasury must model? ● Are AI models that forecast rates, balances, or liquidity governed under model risk?
Third-Party Risk	● Was the appropriate third-party assessment / engagement completed? ● Are there any risks been identified with new third-party relationships? Have they been remediated? ● Are updates on any critical watch statuses within cadence?. ● Is there on going monitoring and coordination activities for vendors with significant third-party risk.	● Is the AI/foundation-model vendor assessed for AI-specific risks — data-usage/training rights, security posture, model transparency/documentation, bias testing, sub-processors, data residency, indemnification, and exit/portability? ● Is there ongoing monitoring of the vendor's model changes/updates that could alter behavior or risk (silent updates, deprecations)? ● Fourth-party risk: does the AI vendor itself rely on other model providers or infrastructure, and is that chain understood? ● Are contractual rights in place to audit, test, and receive notice of material model changes, and to require deletion of our data? ● Does vendor concentration on a single model provider create resilience risk?
Third-Party Technology Solution Risk	● Will there be custom software development (internal or external) supporting the change? ● Will there be third-party technology-solution hosting required? ● Will internal systems send data to a third party? Will a third party send data to internal systems? ● Would operational failure of the solution result in significant business impact?	● Does in-scope business data flow to a third-party model/agent platform, and would a hallucination or operational failure of the agent cause significant business or customer impact? ● Are agent actions on internal systems constrained (least privilege), logged, and reversible, with a kill switch? ● Is the custom AI/agent code securely developed and reviewed (incl. prompt-injection and insecure-output testing)? ● Are data flows to/from the third-party model contractually protected (no training on our data, security, deletion)? ● Is there a fallback if the AI component degrades or the vendor changes the model?