

WHITEPAPER

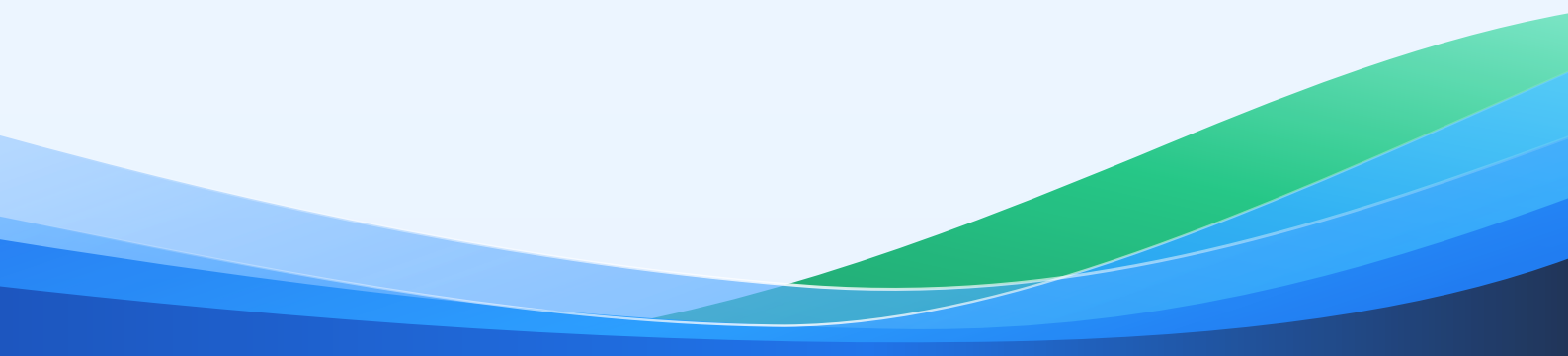
AI in Third-Party Risk Management: Governance, Explainability and Audit Evidence

How to use AI to review vendor evidence without losing accountability, and how to assess the AI your vendors build into their services.

Published 25 September 2026

Written for Heads of TPRM, CROs, CISOs, model risk managers, compliance officers, internal audit

Online www.vendrisk360.com/resources/whitepapers/ai-in-third-party-risk-management/



Contents

01	Why AI in TPRM needs its own governance
02	Where AI helps and where it must not decide
03	Failure modes in AI-assisted evidence review
04	Controls for AI-assisted review
05	Governance frameworks and model risk expectations
06	Explainability and audit evidence
07	How internal audit tests an AI-assisted review
08	Assessing your vendors' use of AI
09	EU AI Act implications
10	Contract clauses for AI
11	How VendRisk360 supports this

A practitioner guide to governing AI inside a TPRM program and assessing vendor AI: controls, failure modes, NIST AI RMF, ISO/IEC 42001, model risk expectations, EU AI Act roles and the evidence auditors will ask for.

AI has arrived in third-party risk management from two directions at once. Internally, TPRM teams are using AI to read the growing stack of SOC reports, policies, penetration test summaries and questionnaires that vendors send them. Externally, the vendors themselves are embedding AI into the services they deliver, often by calling a foundation model provider the customer has never heard of. Both developments raise the same question from boards, examiners and internal audit: who is accountable for the outcome, and what is the evidence?

This paper treats the two sides separately. Part A covers using AI responsibly inside your own TPRM program: where it helps, where it must not decide, how it fails, and the controls and evidence that make it defensible. Part B covers assessing vendors' use of AI: what to ask, what certification does and does not prove, how the EU AI Act allocates responsibility, and which contract clauses matter. The guidance is vendor-neutral and draws on requirements across North America, the EU and UK, India, Singapore, the Middle East and Australia/NZ. Confirm specific obligations with counsel and your regulators.

KEY TAKEAWAYS

- AI can accelerate evidence review (classification, control mapping, SOC exception extraction, draft findings) but must never set risk ratings, accept risk or approve vendors. Those decisions stay with named, accountable people.
- The core control set is simple to state: a named human reviewer for every output, page-level citations that make each claim verifiable, abstention when evidence is missing, a complete audit trail, and change management for models and prompts.
- Treat an AI review tool as an inventoried model or model-like tool. Governance should be proportionate to materiality, which is what supervisors from the US to Singapore expect.
- For document review, "explainable" means traceable: every statement links to a page in the source, and the reviewer's rationale is recorded. That is a different standard from credit model explainability.
- When assessing vendors, ask about AI inventory, training data rights, use of your data for training, foundation model providers as fourth parties, evaluation, prompt injection defenses and human oversight. ISO/IEC 42001 certification is useful evidence of a management system, not proof that a given model is accurate or fair.

Why AI in TPRM needs its own governance

AI changes how the analysis is produced, and that shifts where errors enter. A human analyst who misses a qualified opinion in a SOC 2 report makes a visible, attributable mistake. An AI tool that summarizes the same report and omits the qualification makes an invisible one, and a busy reviewer who trusts the summary inherits it.

The US Interagency Guidance on Third-Party Relationships (June 2023) is clear that the banking organization remains responsible for activities it performs through third parties and for its own risk management decisions. Equivalent principles appear in OSFI B-10, PRA SS2/21, the RBI IT outsourcing Master Direction, MAS guidelines and APRA CPS 230. Nothing in these instruments allows accountability to be delegated to a tool. The practical consequence is that AI output is an input to an accountable decision, never the decision itself, and the program must be able to show that distinction in its records.

Where AI helps and where it must not decide

The dividing line is between tasks that organize or extract information from evidence and tasks that judge what the information means for the organization's risk appetite. The first category benefits from AI speed and consistency. The second requires judgment, context and accountability that belong to people.

AI TASK SUITABILITY IN EVIDENCE REVIEW

Task	AI role	Human role
Document classification (SOC 2, pen test, BCP, policy)	Classify and flag low confidence	Correct misclassifications
Evidence-to-control mapping	Propose mapping with citations	Accept, edit or reject each mapping
SOC report extraction	Extract opinion, exceptions, CUECs, subservice orgs, period	Verify against the report, assess relevance
Drafting findings	Draft wording from cited gaps	Own final wording and severity
Monitoring signal triage	Deduplicate, cluster, summarize	Confirm, dismiss, escalate
Risk rating, risk acceptance, approval	None	Decide and sign off

SOC report extraction deserves specific attention because it is where AI saves the most time and where omissions cost the most. A useful extraction captures the opinion type (unqualified or qualified), the report period and whether it is Type I or Type II, each testing exception with the auditor's description and management's response, the complementary user entity controls (CUECs) the customer must operate, the subservice organizations and whether they are presented using the inclusive or carve-out method, and any complementary subservice organization controls (CSOCs). The reviewer then decides what the exceptions mean for the services actually consumed, whether the CUECs are in place internally, and whether a bridge letter covers the gap between the period end and today.

KEEP THE DECISION BOUNDARY IN THE SYSTEM, NOT JUST THE POLICY

A policy statement that AI does not set ratings is weak evidence if the workflow allows an AI-drafted rating to be approved with one click. Stronger designs make risk rating, acceptance and approval fields human-only, and record who entered each value.

Failure modes in AI-assisted evidence review

Understanding how AI fails in this context lets you design controls against specific errors rather than generic "AI risk". The NIST Generative AI Profile (NIST AI 600-1, July 2024) catalogues risks such as confabulation and information integrity that map directly onto document review.

Hallucination

The tool states something the document does not say: a control that does not exist, a certification the vendor does not hold, a report period that is wrong. The primary defense is requiring a citation for every factual claim and checking that the cited page supports it.

Omission

The tool is accurate in what it says but silent about what matters: the one exception in a 90-page SOC 2 report, a carve-out of the hosting provider, a qualified opinion. Omission is harder to detect than hallucination because there is nothing wrong on the page to catch. Defenses include structured extraction templates that require an explicit value for each field (including "none found"), and QA sampling that re-performs extraction from the source.

Stale or mis-scoped evidence

The document is genuine but does not answer the question: a SOC 2 report whose period ended 14 months ago, a report covering a different legal entity or product line, an ISO/IEC 27001 certificate whose scope excludes the service you buy. AI tools read what they are given. The workflow must check period, entity and scope before any mapping is trusted.

Prompt injection in vendor-supplied documents

Vendor documents are untrusted input. A document can contain text, visible or hidden, that instructs an AI system to mark all controls as satisfied or omit certain content. Defenses include treating document content strictly as data, constraining outputs to structured fields with citations, detecting instruction-like content, and ensuring no AI output changes a risk decision without human review.

Overreliance and automation bias

Reviewers who see accurate outputs most of the time stop checking. Acceptance rates drift toward 100 percent and reviewer time per item drops, which may signal efficiency or may signal rubber-stamping. This is a human factors risk, and it is the reason performance monitoring must include reviewer behavior, not just model accuracy.



Controls for AI-assisted review

The controls below are designed to be testable. Each produces evidence that an auditor can inspect without having to trust the tool.

Human in the loop with a named reviewer

Every AI output that influences an assessment (a classification, a mapping, an extracted exception, a draft finding) is accepted, edited or rejected by a named individual, with timestamp. Group accounts and bulk acceptance undermine this control. Where volume is high, bulk actions should be limited to low-consequence outputs such as document classification and never extend to findings or control conclusions.

Citations and verifiability

Each factual statement carries a page-level citation to the source document. Citations should be clickable in the review interface so that verification takes seconds. A statement without a citation is treated as unsupported.

Confidence and abstention

The tool should be able to say "not found" or "insufficient evidence" rather than guessing. Low-confidence outputs are flagged for closer review. An AI that always produces an answer is more dangerous than one that sometimes declines.

Logging and audit trail

Log the input document version (with a hash), the model and prompt version used, the raw output, the reviewer action, any edits and the final value. This lets you reconstruct any assessment months later, which is exactly what examiners will ask you to do.

Data handling

Vendor evidence often contains confidential and personal data. Controls include: no training of models on customer data unless contractually agreed, tenant isolation so one customer's documents cannot influence or appear in another's outputs, encryption in transit and at rest,

retention limits for prompts and outputs at any model provider, and clarity about where processing occurs (relevant to GDPR, the DPDP Act 2023, PDPA, Saudi PDPL and UAE PDPL).

Change management for models and prompts

A model upgrade or prompt change can alter outputs as much as a code change. Treat them as changes: document the change, run a regression test against a fixed benchmark set of documents with known correct answers, compare results, approve, and record the version in the audit log.

Performance monitoring

Sample reviewed outputs on a regular cadence and classify errors using a stable taxonomy: hallucination, omission, wrong citation, mis-scope, misclassification. Track reviewer override rates by output type and by reviewer. A sudden drop in override rate after a model change, or an individual reviewer who never overrides, both warrant a look.

AI-ASSISTED REVIEW CONTROL MATRIX

Risk	Control	Evidence and test
Hallucinated control or claim	Mandatory page-level citation; reviewer verifies	Sample outputs; confirm cited page supports each claim
Omitted SOC exception or carve-out	Structured extraction with required fields; QA re-performance	Re-extract sample reports; compare to AI output
Stale or mis-scoped evidence	Period, entity and scope checks before mapping	Inspect checks for sample; test with expired report
Prompt injection in documents	Content treated as data; structured outputs; human review	Run seeded test documents; review detection logs
Automation bias	Named reviewer; override and time metrics	Analyze override rates by reviewer and period
Unapproved model or prompt change	Change management with regression benchmark	Trace versions in logs to approved change records
Data leakage or training on customer data	Contract terms; tenant isolation; retention limits	Review provider terms, isolation testing, config
AI decides risk outcome	Human-only rating, acceptance and approval fields	Inspect workflow configuration and audit trail

Governance frameworks and model risk expectations

No single framework is mandatory for AI used in TPRM, but several provide a structure that examiners and auditors recognize.

NIST AI RMF 1.0 and the Generative AI Profile

The NIST AI Risk Management Framework (January 2023) organizes activities into four functions. Govern establishes policies, roles and accountability. Map establishes context: intended use, users, impacts and limitations. Measure assesses and tracks risks through testing and metrics. Manage prioritizes and acts on risks, including decisions to deploy, monitor or retire. The Generative AI Profile (NIST AI 600-1) adds risk categories and suggested actions specific to generative systems. For an AI evidence review tool, Map documents that the tool is decision support for analysts, Measure is the benchmark testing and QA sampling, and Manage is the change management and monitoring described above.

ISO/IEC 42001 and ISO/IEC 23894

ISO/IEC 42001:2023 specifies requirements for an AI management system (AIMS), structured like ISO/IEC 27001: context, leadership, planning (including AI risk assessment and AI impact assessment), support, operation, performance evaluation and improvement, with an annex of AI-specific controls. ISO/IEC 23894 provides guidance on AI risk management, building on ISO 31000. Organizations can align their internal AI governance to these standards without seeking certification.

Model risk management expectations

In the US, SR 11-7 and OCC Bulletin 2011-12 define a model broadly as a quantitative method that processes inputs into estimates, and they expect sound development, independent validation and governance proportionate to use. The PRA's SS1/23 sets model risk management principles for UK banks. In Canada, OSFI's revised Guideline E-23 on model risk management was finalized in 2025 and extends explicitly to AI and machine learning models; check the effective date that applies to your institution. In Singapore, the MAS FEAT principles (fairness, ethics, accountability and transparency) have been in place since 2018, MAS published an information paper on AI model risk management in 2024, and MAS has consulted on guidelines on AI risk management for financial institutions; check the current status. EU and UK supervisors have similarly signaled that AI governance should be proportionate to materiality and embedded in existing risk frameworks.

Is the AI tool a "model"?

A reasonable position for many institutions is that an AI evidence review tool is a model or model-like tool that belongs in the model inventory, even if its risk tier is low because a human verifies every output. Inventorying it costs little and answers the first examiner question. The inventory entry should record intended use, owner, vendor and underlying model provider, materiality tier, validation or testing approach, known limitations and the date of the last review. Where the tool is purchased, the vendor's testing evidence supports but does not replace your own validation of performance on your documents.

Explainability and audit evidence

Explainability means different things for different AI uses, and conflating them leads to either impossible demands or weak evidence.

For a credit scoring model, explainability concerns why the model produced a score: which features drove the outcome, whether reasons can be given to an applicant, and whether the model behaves consistently across groups. That requires techniques that explain the model's internal logic.

For AI-assisted document review, the relevant question is not how the language model works internally but whether each output can be traced and verified. A traceable output has three properties: every claim cites a specific page in a specific document version; the reviewer's action and rationale are recorded; and the final assessment conclusion is written and signed by a person. An auditor can then re-perform the review from the source documents without needing to understand the model at all.

WHAT AUDITORS AND EXAMINERS WILL ASK

- Where is AI used in the third-party risk process, and is it in the model or AI inventory?
- What decisions can the AI make on its own? Show us the workflow configuration that enforces this.
- For this vendor assessment, show every AI output, the reviewer who acted on it and what they changed.
- How do you know the AI did not miss an exception in this SOC report?
- What model and prompt versions were in use on this date, and how were changes approved?
- What testing was done before deployment and after the last model change?
- What are the error rates and override rates, and who reviews them?
- Is vendor data used to train the model? Show us the contract terms and data flow.
- How do you prevent instructions embedded in vendor documents from affecting outputs?

How internal audit tests an AI-assisted review

Internal audit should test the AI-assisted process in the same two stages it applies to any key control: design and operating effectiveness.

A walkthrough follows one assessment end to end: the intake of vendor documents, AI classification and extraction, reviewer actions, draft findings, final rating and sign-off. The auditor confirms that the workflow enforces human-only decision fields, that citations resolve to the

correct pages and that the audit trail captures model version and reviewer identity. Design testing also covers the governance artifacts: inventory entry, testing documentation, change records and the monitoring report.

Operating effectiveness testing typically includes:

1. Re-performance on a sample. Select a risk-based sample of completed assessments, weighted toward critical vendors. Independently extract key facts (opinion, exceptions, CUECs, subservice organizations, period) from the source and compare to the recorded outputs after human review.
2. Citation accuracy check. For a sample of AI statements, open the cited page and confirm it supports the claim. Record mismatches by error type.
3. Reviewer override analysis. Analyze acceptance, edit and rejection rates by reviewer and period. Investigate reviewers with near-zero overrides and any unexplained shifts.
4. Change management trace. Match model and prompt versions in the logs to approved change records and regression test results.
5. Decision boundary test. Confirm no risk rating, acceptance or approval in the sample was recorded by a system account or without a named approver.



HYPOTHETICAL EXAMPLE: RE-PERFORMANCE FINDS A PATTERN OF OMISSION

An internal audit team samples 25 SOC 2 reviews from the prior year. In 23, the recorded exceptions match independent re-performance. In two, the AI extraction listed no subservice organizations where the report used the carve-out method for a cloud hosting provider, and the reviewers accepted the output. Both reports were long, and the carve-out was described in the system description section rather than a dedicated heading. The finding leads to a required "subservice organizations and method" field with mandatory citation, a benchmark test case for this pattern, and a reviewer checklist item. These numbers are illustrative only.

Assessing your vendors' use of AI

Vendors increasingly use AI in the services they provide, from customer support to fraud detection. A vendor's AI can process your data, affect your customers and introduce new fourth parties. Add an AI domain to your questionnaires and evidence requests for any vendor whose service uses AI on your data or in a decision affecting your customers.

VENDOR AI ASSESSMENT DOMAIN

Area	What to ask	Evidence to request
AI inventory and intended use	Which AI systems touch our service or data? For what purpose?	AI inventory extract; system descriptions
Training data provenance	Where did training data come from? Are rights documented?	Data provenance summary; licensing position
Customer data use	Is our data used to train or improve models? Can we opt out?	Contract terms; data flow diagram; configuration
Foundation model providers	Which model providers and hosting regions are used?	Subprocessor list; provider terms on retention and training
Evaluation and bias testing	How is accuracy measured? Is fairness tested where relevant?	Test methodology; recent results summary
Security	How are prompt injection, data leakage and model abuse addressed?	Threat model; pen test scope including AI features
Human oversight	Which outputs are reviewed by people before use?	Process documentation; oversight roles
Incident handling	Are AI failures and misuse covered by incident response?	IR plan; notification commitments
Governance and certification	Is there an AI management system? Certified to ISO/IEC 42001?	Certificate and scope statement; AI policy

Foundation model providers are fourth parties

When a vendor builds a feature on a third-party foundation model accessed through an API, that model provider becomes your fourth party. Several regimes make this visible: DORA requires attention to subcontracting of ICT services supporting critical or important functions, GDPR Article 28 requires prior authorization of sub-processors, APRA CPS 230 expects entities to consider fourth parties that material service providers rely on, and PRA SS2/21 addresses sub-outsourcing. Map foundation model providers in your nth-party inventory. Because a handful of providers supply a large share of the market, the same provider may sit behind many of your vendors, creating concentration risk that is invisible if you look at vendors one at a time.

What ISO/IEC 42001 certification does and does not prove

A certificate shows that an accredited body audited the vendor's AI management system against the standard and found it conforming within a defined scope. It is useful evidence that governance, risk assessment, impact assessment and lifecycle processes exist. It does not prove that a particular model is accurate, unbiased or secure, that the certified scope covers the service you buy, or that the vendor meets any specific legal obligation such as the EU AI Act. Always check the scope statement, the certification body and its accreditation, and the certificate dates, and pair the certificate with evidence about the specific AI features in your service.



EU AI Act implications

The EU AI Act, Regulation (EU) 2024/1689, entered into force on 1 August 2024 and applies in phases. It is relevant to TPRM because it assigns obligations by role, and your organization and your vendors may hold different roles for the same system. It can apply to organizations outside the EU where AI outputs are used in the EU.

Provider versus deployer

A provider develops an AI system (or has it developed) and places it on the market or puts it into service under its own name. A deployer uses an AI system under its authority in a professional context. A bank that buys a vendor's credit scoring system is typically a deployer; the vendor is typically the provider. Substantial modification or rebranding can shift an organization into the provider role, so check this carefully when customizing vendor AI.

High-risk uses in financial services and insurance

Annex III lists high-risk uses that include AI systems used to evaluate the creditworthiness of natural persons or establish their credit score (with an exception for fraud detection) and AI systems used for risk assessment and pricing in relation to natural persons for life and health insurance. Other Annex III categories, such as certain employment uses, can also apply to financial institutions as employers.

Deployer obligations

For high-risk systems, deployers are expected, among other things, to use the system in accordance with the provider's instructions, assign human oversight to competent people, ensure input data under their control is relevant, monitor operation and inform the provider of risks or serious incidents, keep automatically generated logs under their control for an appropriate period, and inform affected individuals where required. Certain deployers, including those evaluating creditworthiness or pricing life and health insurance, must carry out a fundamental rights impact assessment. Check the current text for the precise scope of each obligation.

Timetable

Prohibited practices applied from February 2025 and general-purpose AI model obligations from August 2025. Most high-risk obligations were scheduled to apply from August 2026, but amendments proposed by the European Commission would link or defer some of these dates. Check the current timetable before relying on any date.

For TPRM, the practical steps are to identify vendor AI systems that may fall into Annex III uses, confirm who is provider and who is deployer, obtain the provider's instructions for use and conformity information, and make sure contracts support your deployer obligations.

Contract clauses for AI

Standard technology contracts rarely address AI specifically. The clauses below can be added to master agreements, data processing agreements or AI-specific addenda, scaled to the materiality of the AI use.

AI CONTRACT CLAUSES

Clause	Purpose
AI use disclosure	Vendor discloses AI features that process customer data or affect decisions, and notifies material changes
No training on customer data	Customer data, prompts and outputs are not used to train or improve models without written agreement
Model provider flow-down	Foundation model and hosting providers listed as subprocessors, with equivalent obligations and prior notice of changes
Data residency and retention	Where prompts and outputs are processed and how long any provider retains them
Evaluation and transparency	Vendor provides testing results, known limitations and instructions for use on request
Human oversight support	Vendor supplies features and information needed for customer oversight and logging
Security of AI features	Prompt injection, data leakage and abuse are in scope for security testing and remediation
AI incident notification	Timely notice of AI malfunctions, harmful outputs or data exposure, aligned to your incident regime
Regulatory cooperation	Vendor supports customer obligations under applicable AI laws, including the EU AI Act where relevant
Audit and information rights	Customer and regulators can obtain evidence about AI features, consistent with existing audit clauses
IP and output rights	Ownership of outputs and indemnity for third-party IP claims arising from model use

- ☐ AI domain added to questionnaires for vendors using AI on your data or decisions
- ☐ Foundation model providers recorded as fourth parties in the nth-party inventory
- ☐ ISO/IEC 42001 certificates checked for scope, body and dates
- ☐ EU AI Act role and Annex III relevance recorded for each vendor AI system in scope
- ☐ AI clauses included in new and renewing contracts for material AI use
- ☐ Your own AI review tool inventoried, tested and monitored

How VendRisk360 supports this

VendRisk360 applies the principle this paper describes: AI is scoped narrowly, and people do the review and make the decisions.

For Part A, AI is an optional capability that each customer chooses whether to switch on. Where a customer opts in, AI assists in exactly two places: **completeness checks and key-date extraction** on vendor evidence (effective and expiration dates, period covered, issuer, document type and scope, with missing, expired or out-of-scope items flagged), and an **AI-assisted first pass on SOC reports** within [SOC Report Review](#), which the expert assessor verifies. Every review, rating and sign-off is performed by a VendRisk360 expert assessor or by the customer's own reviewers, and the platform workflow does not depend on AI. Where AI is used, each output is accepted, edited or rejected by a named person and the action is recorded in the audit trail, which supports the re-performance and override analysis described above. **SOC 1 and SOC 2 review** covers exceptions, carve-outs, subservice organizations and CUECs. AI never sets risk ratings, accepts risk or approves vendors: findings, remediation plans and formal risk acceptance run through **multi-stage sign-off with segregation of duties** and a sign-off certificate. Our AI governance is aligned to the NIST AI RMF and ISO/IEC 42001; VendRisk360 is not certified to ISO/IEC 42001. Customer data is protected by **row-level-security tenant isolation**, encryption in transit and at rest, SSO (SAML/OIDC), mandatory MFA, custom roles and a full audit trail. See [Security](#).

For Part B, **artifact-based assessments across 30+ control domains** and questionnaires let you add AI-specific evidence requests to vendors whose services use AI, scoped by criticality tier, with reassessment cadence configured to your own policy. The vendor portal gives vendors one-time-code access to only their own requests. **Nth-party intelligence** maps fourth-party and nth-party relationships across your vendors, so a foundation model provider behind several vendors surfaces as a concentration and blast radius question rather than a hidden dependency. See [Nth-party intelligence](#) and [Comprehensive Vendor Risk Assessment Services](#).

When examiners ask to see how an assessment was reached, the **examiner package export** bundles assessments, evidence, sign-offs and the audit trail, and the Audit Committee / Examiner Readiness deck summarizes program status for oversight bodies. See [Board and executive reporting](#).

How to work with VendRisk360. On the [Vendor Lifecycle Management Platform](#), your team manages vendors, sends due diligence and evidence requests through the vendor portal, performs the review, records it and signs off, with every step tracked in a full audit trail. With [Comprehensive Vendor Risk Assessment Services](#), you onboard the vendor and VendRisk360's certified assessors collect the evidence and follow up with the vendor, perform the risk assessment scaled to the vendor's tier with a second-expert quality review, and follow findings through remediation. You see each vendor's progress on the platform in near real time and keep final approval, as regulators expect. [Continuous Monitoring Services](#) add the outside-in view between assessments, and [Report-Specific Reviews](#) (SOC Report Review, Information Security Program Review and Business Continuity Program Review) are available standalone or alongside either. Assessors hold

certifications such as CISSP, CISA, CISM, CRISC and ISO/IEC 27001 Lead Auditor and Lead Implementer, with PCI DSS implementation experience. AI is optional: where a customer opts in, it assists only with completeness checks and key-date extraction on vendor evidence and with an AI-assisted first pass on SOC reports that the expert assessor verifies. Every review, rating and sign-off is performed by an expert assessor or by the customer's own reviewers.

To see SOC Report Review with optional AI assistance, and human sign-off, in practice, [book a demo](#).



About VendRisk360

VendRisk360 is an independent company providing a third-party risk management platform and expert services for regulated organizations: banks, credit unions, fintech and payments, healthcare and SaaS. The platform runs vendor lifecycle management, risk-tiered assessments, continuous monitoring, nth-party intelligence and board-ready reporting from one governed record per vendor, combining point-in-time, evidence-based review with continuous outside-in monitoring. Customers run it themselves on the Vendor Lifecycle Management Platform, or add Comprehensive Vendor Risk Assessment Services, Continuous Monitoring Services and Report-Specific Reviews delivered by certified VendRisk360 assessors.

Learn more at vendrisk360.com or write to info@vendrisk360.com.

This paper is general guidance for practitioners, not legal advice. Regulatory requirements change and vary by jurisdiction and institution; confirm your obligations with counsel and your supervisors.