

AI-Driven Due Diligence: Compressing CDD, FDD and TDD Without Losing Rigour

An assertion-led evidence architecture for faster review, deterministic analysis and accountable transaction judgement

Chennakeshav (CK) Adya

Independent Researcher

July 2026

ABSTRACT

Background. Artificial intelligence is entering M&A workflows, including diligence and valuation, while transaction evidence remains fragmented, time-constrained and judgement-intensive. **Objective.** This paper develops a controlled operating model for AI-assisted commercial, financial and technology due diligence. **Approach.** The analysis reviews 18 primary, authoritative and practitioner sources. Survey findings retain their sample and self-report boundaries. Selected PCAOB audit-evidence principles are used as a clearly limited analogy for financial due diligence. **Findings.** KPMG reports that 56% of 700 surveyed dealmakers deployed agentic AI in due diligence and valuation. Thomson Reuters reports organisation-wide GenAI adoption among 40% of 1,514 surveyed professionals, while 18% said their organisations collected AI ROI metrics. The proposed architecture converts the mandate into atomic assertions; registers every source and version; uses AI for controlled extraction, comparison, triage and drafting; retains deterministic financial calculations and reproducible technical tests; and routes material exceptions to named reviewers. **Implications.** Time compression and quality preservation are pilot questions. Source coverage, citation validity, calculation fidelity, exception performance, reviewer time, elapsed time and rework provide a measurable release record.

Highlights

- The assertion, rather than the document, is the operating unit of controlled diligence.
- AI can expand ingestion, comparison, exception triage and cited drafting capacity.
- Financial calculations remain deterministic and technical tests remain reproducible.
- Material contradictions, missing evidence and failed checks enter an owned exception queue.
- Compression and rigour require a defined baseline and observed pilot evidence.

JEL Classification: G24, G34, M15, M41, O33

Keywords: artificial intelligence, due diligence, commercial due diligence, financial due diligence, technology due diligence, M&A, private equity, audit evidence, evidence graph, model risk, transaction execution

Research paper only. This document is not legal, tax, accounting, cyber-security or investment advice. Matchpoint scenario analysis is illustrative and is not observed productivity, a forecast or evidence of improved transaction returns.

1. INTRODUCTION

Due diligence is an exercise in controlled doubt. An investment committee asks whether the business, earnings and technology described by a seller can withstand scrutiny before capital is committed. The answer is assembled from imperfect records, compressed access, changing data-room populations, management explanations and specialist judgement. Artificial intelligence adds a new analytical layer to that process. It can read, classify and compare large collections of documents; propose mappings across inconsistent data; draft issue summaries; and direct reviewers towards exceptions. It can also produce confident statements that are unsupported, conflate periods or entities, overlook a missing population and make a calculation appear coherent when its inputs are wrong.

Current market evidence shows that AI has entered transaction work. KPMG's 2026 survey covered 700 senior M&A decision-makers across 20 countries and jurisdictions. Fifty-six per cent reported deploying agentic AI in due diligence and valuation. In the same survey, 59 per cent reported an efficiency gain above 10 per cent in competitive intelligence and market analysis, while 17 per cent reported an efficiency gain above 25 per cent in valuation modelling and scenario planning and in budget modelling and tracking [1,2]. These are respondent-reported findings. They do not establish causal productivity, error reduction or investment performance.

Measurement is materially less mature than adoption. The Thomson Reuters Institute surveyed 1,514 professionals across legal; tax, audit and accounting; corporate risk and fraud; and government functions in October and November 2025. Forty per cent reported organisation-wide generative-AI use, compared with 22 per cent in the prior edition. Eighteen per cent reported that their organisations collected AI return-on-investment metrics, and 40 per cent did not know whether such metrics were collected [3]. The sample extends beyond transactions. It provides evidence of a broad professional-services measurement gap rather than an M&A-specific outcome.

This paper develops an operating model for AI-assisted commercial due diligence, financial due diligence and technology due diligence, abbreviated CDD, FDD and TDD. The model begins with the assertion to be tested. Every conclusion is connected to a source, version, transformation, calculation or test, exception and reviewer. The model separates generative work from deterministic calculation and reproducible testing. It assigns named people to resolve material exceptions and authorise conclusions.

The paper uses selected PCAOB audit-evidence principles as an analytical analogy for FDD. PCAOB standards govern audits performed within their scope; they do not govern every financial due-diligence mandate. The analogy is useful because it distinguishes the quantity of evidence from its relevance and reliability, treats contradictory material as evidence, links the persuasiveness of evidence to risk and differentiates inquiry from recalculation and reperformance [4-6]. The transaction work programme and contractual duty remain engagement-specific.

1.1 Research questions

The analysis addresses five questions:

- Which CDD, FDD and TDD tasks are suitable for AI-assisted ingestion, extraction, comparison and drafting?
- Which calculations, tests, judgements and authorisations should remain deterministic or specialist-controlled?
- How can a diligence team trace each material conclusion back to evidence and forward from a changed source to affected conclusions?
- Which measures distinguish elapsed-time reduction from reviewer-time reduction and coverage expansion?
- Which governance, security and decision-rights controls support a bounded production deployment?

1.2 Executive conclusion

[Matchpoint synthesis] AI changes diligence capacity more readily than it changes the evidence threshold. The immediate opportunity is parallel work: controlled ingestion, first-pass classification, population-wide comparison, exception ranking and cited drafting can operate while specialists investigate the most material questions. The appropriate operating unit is an assertion, not a document. A document can support several assertions and a single assertion can require many documents, calculations, tests and interviews.

[Matchpoint synthesis] A controlled implementation contains six parts. First, the mandate is translated into atomic assertions with a materiality and evidence rule. Second, every source is registered with identity, version, scope and access. Third, extraction and classification create traceable proposed facts rather than final conclusions. Fourth, deterministic engines perform arithmetic and reproducible tools perform technical tests. Fifth, exceptions are routed to an owner under a defined closure policy. Sixth, an authorised reviewer signs the conclusion and the decision record preserves the basis and limitations.

The words "without losing rigour" describe a control objective with no assured performance result. A pilot should compare the assisted workflow with a defined baseline using source coverage, citation validity, calculation fidelity, unresolved exceptions, reviewer time, elapsed time, rework and issue yield. Quality claims require observed results from that pilot. Transaction returns remain outside the evidence available for this paper.

Three disciplines, one controlled operating model



Figure 1. CDD, FDD and TDD retain specialist domains and share an evidence, exception and sign-off system

Source: Matchpoint framework.

2. ADOPTION, EFFICIENCY AND THE MEASUREMENT GAP

2.1 What current surveys establish

KPMG's 2026 study is directly relevant because its respondents were senior transaction decision-makers. The survey reports deployment across diligence, valuation, sourcing, execution, integration and compliance. It also describes a shift in economic scope: lower marginal analysis cost can make broader contract review, competitive benchmarking and historical pattern analysis more feasible [1]. This is practitioner and survey evidence. It supports a hypothesis about expanded coverage. A production team still needs to measure the actual population processed, the exceptions reviewed and the accuracy achieved.

The survey's adoption figure should be read with its method. Respondents numbered 700; 519 were corporate and 181 were from private-equity firms. Fieldwork ran from 19 December 2025 to 27 January 2026. The survey covered 20 countries and jurisdictions and 10 sectors. Respondents held responsibility for, or involvement in, M&A decisions [1]. The reported 56 per cent therefore describes that sample and the survey's definition of agentic-AI deployment. It does not establish that 56 per cent of all deal teams use a comparable workflow.

The reported efficiency findings are useful as a market signal. They are not a transaction-control certificate. The press release states that 59 per cent reported more than 10 per cent efficiency gain in competitive intelligence and market analysis. Seventeen per cent reported more than 25 per cent efficiency gain in valuation modelling and scenario planning and in budget modelling and tracking [2]. The disclosed material

does not provide an independently observed baseline for each respondent, a standard work definition or a verified quality-adjusted output measure. The paper therefore uses the figures to motivate measurement, not to set Matchpoint targets.

Thomson Reuters provides a wider professional-services perspective. Its survey was online, included 1,514 respondents across 27 countries and screened participants for familiarity with AI. The majority of respondents came from the United States, United Kingdom and Canada. Common applications included research, drafting and summarisation [3]. These tasks resemble parts of the diligence production chain, particularly initial document review and report assembly. The report's ROI findings indicate that formal measurement may lag day-to-day use.

Evidence	Sample and period	Reported result	Interpretation boundary
KPMG M&A survey [1,2]	700 senior dealmakers; 20 countries and jurisdictions; Dec 2025-Jan 2026	56% reported agentic AI in diligence and valuation	Self-reported deployment; no universal quality or outcome claim
KPMG M&A survey [2]	Same sample	59% reported more than 10% efficiency gain in competitive intelligence and market analysis	Metric and baseline vary by respondent
Thomson Reuters [3]	1,514 professionals; 27 countries; Oct-Nov 2025	40% reported organisation-wide GenAI use; 18% reported ROI measurement	Broad professional-services sample; transaction-specific inference is unavailable

2.2 Four different forms of speed

[Matchpoint synthesis] Diligence teams should distinguish four quantities. Machine processing time measures the time used to ingest or analyse a defined input. Reviewer time records the human effort needed to configure, inspect, correct and approve the output. Elapsed time measures calendar time from a defined start event to a defined completion event. Coverage expansion records additional sources, records or tests completed within a comparable budget and timetable.

A workflow can reduce machine processing time and increase reviewer time if outputs create many low-value exceptions. It can reduce reviewer time while leaving elapsed time unchanged because management responses or external confirmations control the critical path. It can preserve elapsed time while expanding population coverage. Each result is economically different. A pilot scorecard should record them separately.

2.3 The measurement problem

The Thomson Reuters findings place the measurement question at the centre of implementation. Where only 18 per cent of respondents report ROI collection and 40 per cent do not know whether it occurs, adoption statistics provide limited evidence about business value [3]. A transaction team needs measures attached to the work programme itself. Generic employee-use counts or tokens consumed do not establish diligence quality.

[Matchpoint synthesis] The primary unit of measurement is the assertion-procedure pair. For each material assertion, the team records expected evidence, sources received, procedures completed, exceptions raised, exceptions resolved, reviewer minutes, completion date and conclusion status. Aggregate programme measures are constructed from these records. This permits a reviewer to identify whether a faster report resulted from genuine workflow improvement, reduced scope or a lower evidence threshold.

3. WHAT RIGOUR MEANS IN TRANSACTION DILIGENCE

3.1 Sufficiency, relevance and reliability

PCAOB AS 1105 describes sufficiency as the quantity of audit evidence and appropriateness as its relevance and reliability [4]. The standard also states that evidence includes information that supports and corroborates

management assertions and information that contradicts them. These are audit requirements within PCAOB scope. [Matchpoint synthesis] The same distinctions form a useful diligence design discipline.

A large data room can be insufficient for a material question. Ten versions of an internally prepared market forecast may have less evidential value than a smaller set of independent, current and perimeter-matched sources. A ledger export can be complete as a file and incomplete as a population if entities, periods or journals are omitted. A repository scan can process every file and remain irrelevant to production architecture if the supplied revision differs from the deployed code.

[Matchpoint synthesis] Each assertion should therefore define an evidence threshold along four dimensions:

- **Coverage:** which entities, periods, customers, products, repositories or environments are in scope.
- **Relevance:** how closely the source and procedure address the assertion.
- **Reliability:** source independence, controls, originality, lineage and susceptibility to alteration.
- **Recency:** the date at which the evidence is valid and the event that causes it to expire.

3.2 Risk-weighted evidence

PCAOB AS 2301 links higher assessed risk to more persuasive evidence and requires procedures that address relevant assertions [5]. It also describes professional scepticism as a questioning mind and critical assessment of evidence. [Matchpoint synthesis] A diligence workflow can adapt this principle through an assertion risk score that governs minimum evidence and reviewer seniority.

High-materiality assertions may require an independent source, deterministic recalculation, specialist review or direct test. Lower-materiality assertions may close with a documented internal source and plausibility check. The score does not automate judgement. It makes the intended level of work visible and allows departures to be approved.

3.3 Inquiry, extraction, recalculation and reperformance

An interview statement is inquiry. An AI summary of that interview remains a representation derived from inquiry. It does not become corroboration through fluent drafting. AS 1105 states that inquiry alone is insufficient to support an audit conclusion within the standard's scope [4]. A transaction mandate may apply a different threshold. The underlying distinction remains valuable.

Recalculation checks mathematical accuracy. Reperformance independently executes procedures or controls [4]. [Matchpoint synthesis] FDD should use deterministic recalculation for material bridges and ratios. TDD should use reproducible tests where access and mandate permit. CDD may use independent market triangulation, customer data analysis, contract evidence and interviews across different constituencies.

Evidence procedure	What it establishes	AI role	Closure authority
Inquiry	A person's representation at a stated time	Transcribe, code themes, identify follow-ups	Reviewer records representation and corroboration status
Inspection	Content of a source object	Extract, classify, compare, cite	Reviewer confirms source, scope and material terms
Recalculation	Mathematical accuracy under defined inputs and rules	Propose mapping and explain output	Deterministic engine calculates; specialist approves method
Reperformance	Independent execution of a procedure or control	Prepare test plan and triage results	Specialist executes or validates reproducible result
External corroboration	Evidence from an independent source	Discover and align candidates	Reviewer establishes identity, relevance and date

3.4 The evidence graph

[Matchpoint synthesis] A diligence report should operate as a view over an evidence graph. The graph contains nodes for assertions, sources, source versions, extracted facts, transformations, calculations, tests, exceptions, conclusions and reviewers. Edges record relationships such as supports, contradicts, derived from, tested by, supersedes, reviewed by and closed by.

This structure supports two forms of traceability. Backward traceability moves from a report sentence to the conclusion, procedure, fact and source. Forward traceability starts with a changed or withdrawn source and identifies every affected fact, calculation, exception and report sentence. Conventional folders and footnotes provide partial backward traceability. A graph adds dependency visibility.

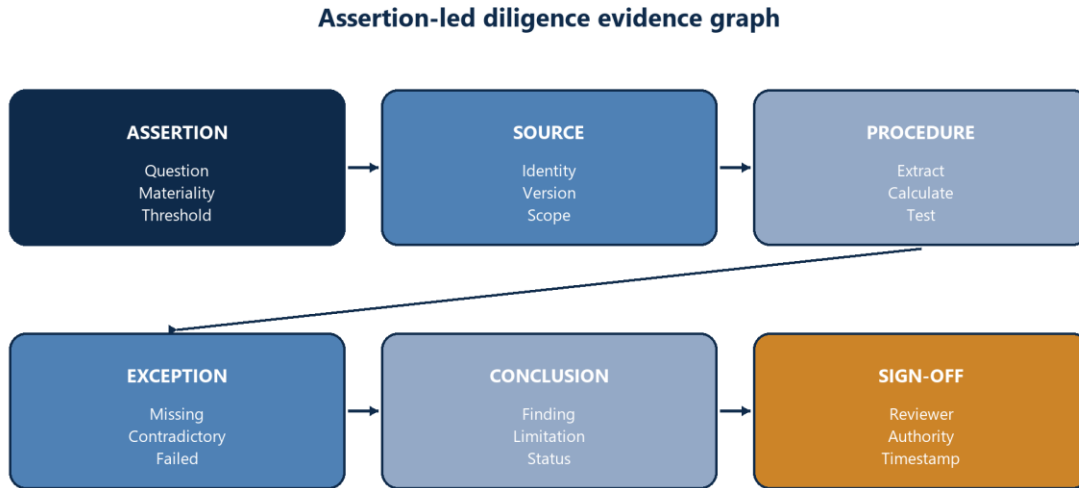


Figure 2. Evidence graph from assertion through source and procedure to exception, conclusion and sign-off

Source: Matchpoint framework, informed by audit-evidence concepts in [4-6] and information-integrity guidance in [8].

4. ONE OPERATING MODEL ACROSS THREE DIFFERENT DISCIPLINES

CDD, FDD and TDD examine different assertion families. They can share an ingestion, evidence, exception and sign-off system. Shared infrastructure reduces duplication while preserving professional boundaries.

4.1 Common sequence

[Matchpoint synthesis] The common sequence has eight stages:

- **Mandate:** define decision questions, scope, materiality, jurisdictions, reliance and exclusions.
- **Assertion map:** convert questions into testable statements with evidence thresholds.
- **Source registration:** hash, identify, date, classify and permission every source object.
- **Controlled extraction:** produce proposed facts with source locations and confidence metadata.
- **Analysis:** run deterministic calculations, reproducible tests and source triangulation.
- **Exception triage:** rank contradictions, missing evidence, variances and failed tests.
- **Specialist review:** validate methods, evidence and judgement under named authority.
- **Decision record:** publish the conclusion, limitations, unresolved issues and sign-off.

4.2 Shared controls

The shared controls include access by role and deal; segregation of confidential workspaces; source hashes and versions; data-loss controls; approved model and tool inventory; prompt and configuration versioning; logs; evaluation sets; exception ownership; retention; incident procedures; and human authorisation. NIST's AI RMF provides a voluntary Govern, Map, Measure and Manage structure for this work [7]. NIST's Generative AI Profile identifies confabulation, privacy, information-integrity, security and human-AI configuration risks that are directly relevant to evidence-heavy work [8].

4.3 Different professional boundaries

The common operating model does not merge professional accountabilities. Commercial specialists decide whether market and customer evidence supports the investment thesis. Accounting and transaction-services specialists decide the treatment of earnings, debt, working capital, tax and accounting matters. Technology and cyber specialists interpret architecture, security, delivery and technical debt. Counsel determines applicable law and legal conclusions. The investment committee owns the investment decision.

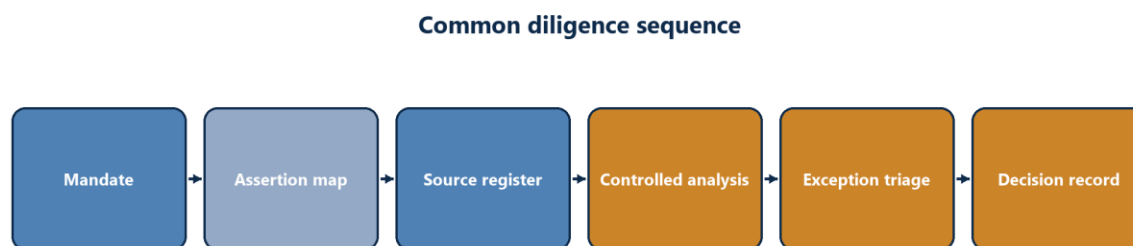


Figure 3. Common sequence from mandate to authorised decision record

Source: Matchpoint framework.

5. COMMERCIAL DUE DILIGENCE

5.1 Assertion families

[Matchpoint synthesis] Commercial assertions usually concern market size and growth; segment definition; customer needs and switching; concentration and retention; product differentiation; pricing power; route to market; competitive response; regulatory exposure; and the credibility of the commercial forecast. Each assertion requires a perimeter. A growth rate without geography, period, segment, currency and nominal or real basis is incomplete evidence.

AI is well suited to source discovery, structured extraction, taxonomy alignment, interview coding, contract-term comparison and first-draft synthesis. It can align statements across market reports and identify where definitions differ. It can cluster customer interview themes and link each theme to the underlying transcript. It can compare product claims with documented features and contract commitments.

5.2 Market triangulation

[Matchpoint synthesis] A market model should retain each source's definition and date. The AI layer can map source taxonomies to the transaction taxonomy and flag unmatched categories. The deterministic layer performs conversions and aggregation. The specialist decides whether a mapping is economically valid.

For example, a target may describe its addressable market as regional enterprise software. External sources may classify revenue by industry, customer size, deployment model or geography. A generated summary can hide these differences. The evidence graph should preserve them. The conclusion may present a range, a scenario or a scope limitation where definitions cannot be reconciled.

5.3 Customer and revenue evidence

Customer-level data can support concentration, retention, cohort, cross-sell and pricing analysis. The AI layer can propose customer-name normalisation, contract-field extraction and theme coding. Deterministic calculations should produce revenue concentration, gross and net retention, cohort bridges and price-volume-mix analysis under documented rules.

[Matchpoint synthesis] Every normalisation proposal should retain original values and the mapping rule. Material customer mergers, channel partners, resellers and pass-through revenue need specialist review. A model that groups similar names can create a plausible but incorrect customer family. The workflow should direct uncertain or high-value mappings to an exception queue.

5.4 Competition and regulation

Competition analysis is jurisdiction-specific. The US DOJ and FTC 2023 Merger Guidelines provide an official US framework and analytical, economic and evidentiary tools [11]. They do not define the legal analysis for every transaction. Counsel should determine relevant jurisdictions, filing regimes and legal conclusions.

[Matchpoint synthesis] AI can assist by building a dated competitor evidence set, comparing product and pricing claims and tracing market-share inputs. It should label source provenance and preserve dissenting evidence. Legal characterisation and privilege decisions remain with counsel.

5.5 CDD output

A controlled CDD output includes the assertion, perimeter, evidence received, method, finding, contrary evidence, sensitivity, limitation and reviewer. The report can then distinguish a verified historical observation, a management estimate, an external forecast and a Matchpoint scenario. This separation is especially important when the investment case depends on future growth.

6. FINANCIAL DUE DILIGENCE

6.1 Data lineage and tie-out

FDD begins with population control. The team should identify entities, ledgers, periods, currencies, account structures, consolidation adjustments and the relationship between management accounts, statutory financial statements and transaction schedules. AI can propose field mappings and identify likely duplicates or missing periods. A deterministic process should count records, total control fields, hash inputs and reconcile outputs.

PCAOB AS 2301 states, within audit scope, that period-end substantive procedures include reconciliation of financial statements to underlying accounting records and examination of material adjustments [5]. [Matchpoint synthesis] A diligence work programme can use this as a design analogy: derived metrics should connect to the supplied accounting population and all material transformations should be visible.

6.2 Quality of earnings

AI can extract management-proposed EBITDA adjustments, classify supporting documents and compare narratives across versions. It can suggest whether an item resembles non-recurring cost, run-rate saving, owner expense, revenue adjustment or accounting-policy difference. The specialist determines the adjustment category, evidential support, period, tax effect and treatment.

[Matchpoint synthesis] The deterministic bridge should begin with a defined reported earnings measure, apply separately identified adjustments, preserve reversals and calculate the adjusted measure. Each adjustment connects to source evidence and reviewer status. Generated prose explains a calculated bridge; it does not calculate the bridge.

6.3 Working capital and net debt

Working-capital and net-debt schedules frequently contain classification judgement, seasonal patterns, cut-off issues and deal-definition questions. AI can extract contractual definitions and propose account mappings. Deterministic code or controlled spreadsheet logic should calculate monthly balances, normalisation ranges, seasonality, cash-like and debt-like items and sensitivity cases.

The purchase agreement and applicable accounting advice govern the transaction definition. A model can compare draft clauses and schedules. Counsel and financial specialists approve their interpretation. IFRS 3 establishes principles for recognition and measurement of acquired assets and liabilities, goodwill or bargain purchase and related disclosures where the standard applies [14]. Transaction accounting remains separate from the diligence model.

6.4 Forecast and cash conversion

Forecast diligence connects historical drivers to management assumptions. AI can map narrative assumptions to model inputs and flag inconsistent descriptions. Deterministic analysis should calculate historical conversion, driver bridges, sensitivities and scenario outputs. The commercial team tests demand and pricing assumptions; the financial team tests margin, working-capital, capital-expenditure and cash consequences.

6.5 Evidence hierarchy and scepticism

[Matchpoint synthesis] A financial evidence hierarchy might rank direct bank or third-party confirmation, controlled system data, original contracts and invoices, management schedules, interviews and generated summaries differently according to the assertion. The hierarchy is engagement-specific. It should be documented before conclusions are formed.

PCAOB guidance states that higher risk requires more persuasive evidence and that inquiry alone does not provide sufficient evidence for the cited audit conclusions [4,5]. In an FDD context, the practical lesson is to preserve the distinction between explanation and corroboration. Management commentary can explain a variance. Source records, calculations or independent evidence may be required to close it under the agreed work programme.

7. TECHNOLOGY AND AI DUE DILIGENCE

7.1 Technology evidence domains

[Matchpoint synthesis] TDD covers product architecture, software repositories, cloud and infrastructure, cyber security, data architecture, engineering and release processes, service resilience, third-party dependencies, intellectual property, licences, technical organisation and cost-to-scale. Where the target develops or deploys AI, the scope expands to models, training and evaluation data, prompts, retrieval sources, evaluation sets, safety controls, monitoring, vendor dependencies and model-change management.

Documentary evidence is necessary and incomplete. Architecture diagrams describe intended structure. Configuration exports and repository evidence show implemented structure. Tests and operational logs show observed behaviour within a time and environment. Interviews explain decisions and known limitations. The report should preserve these evidence classes.

7.2 Secure-development and supply-chain lineage

NIST SP 800-218A supplements secure software-development practices with AI-specific considerations and is intended for model producers, AI system producers and acquirers [9]. The NCSC-led secure-AI guidance organises practices across secure design, development, deployment, and operation and maintenance. It also addresses supply chain, documentation, models and data, logging and incident management [10].

[Matchpoint synthesis] An AI-target diligence record should identify the model or service, provider, version, intended use, system components, training or fine-tuning status, evaluation evidence, retrieval sources, data

rights, deployment environment, access, monitoring, incidents and change process. Where access permits, the team should link this inventory to repositories, manifests, model cards, bills of materials, configuration, tests and logs.

7.3 Reproducible tests

AI can propose test cases, explain results and rank findings. A reproducible test record should preserve the target revision, environment, command or procedure, tool version, inputs, timestamp, raw output, severity rule and reviewer. Reproduction may be constrained by access, production safety, data sensitivity or time. The report should state those constraints.

A vulnerability scanner finding does not by itself establish exploitability or business impact. A clean scan does not establish the absence of vulnerabilities. Specialist review connects the observed result to architecture, exposure, compensating controls and transaction materiality.

7.4 Data, privacy and cyber disclosures

NIST's Generative AI Profile identifies privacy and information-security risks, including leakage, inference and malicious interaction with AI systems [8]. The NCSC guidance recommends secure treatment of models, data, infrastructure and logs [10]. The SEC's current small-entity compliance guide summarises US public-company cyber incident and risk-management, strategy and governance disclosure requirements [15]. Its relevance depends on issuer status and jurisdiction.

The UAE federal personal-data framework may be relevant to a UAE transaction or data room. The official legislation page was access-restricted to the research browser during this study, so this paper makes no detailed legal proposition from its text [13]. Qualified counsel should determine federal, free-zone, sector and cross-border requirements.

7.5 AI exposure as an investment question

Technology diligence now includes the target's exposure to AI as a competitive force and an operational dependency. PwC's analysis of the 100 largest corporate M&A deals in 2025 reports that approximately one-third cited AI in the strategic rationale [17]. This selected sample does not represent the whole market. It indicates that AI can be both an execution tool and a subject of diligence.

[Matchpoint synthesis] The investment question has at least four parts: whether AI changes customer willingness to pay; whether it reduces entry barriers or differentiation; whether the target depends on concentrated model or cloud suppliers; and whether the target has rights, controls and evidence for its own AI system. These questions require commercial, technical and legal collaboration.

8. CONTROLLED ARCHITECTURE

8.1 Transaction data plane

[Matchpoint synthesis] The transaction data plane stores original source objects and registered derivatives inside a deal-specific security boundary. Each object receives an identifier, cryptographic hash, source, received date, document date, entity, period, confidentiality class and access rule. New versions do not overwrite old versions. The system records supersession and the conclusions affected.

Sensitive personal, customer, employee, pricing and technical data should only enter approved systems under the mandate and applicable law. Public consumer tools may be outside the approved boundary. The implementation should document model-provider retention, training use, region, subprocessors and access. Counsel and information-security owners determine acceptable configuration.

8.2 Ingestion and retrieval

Files are scanned, parsed and indexed under an allow-list. The ingestion pipeline should preserve page, cell, record or code-location references. Table extraction needs specific validation because visual structure can be

lost. Prompt injection and malicious embedded content are treated as input threats. Retrieval returns only sources authorised for the user and transaction.

[Matchpoint synthesis] The generated answer should carry evidence pointers at the claim level. A citation to an entire 200-page document may be too coarse for a material conclusion. The interface should allow a reviewer to open the exact passage, table row, calculation input or test record.

8.3 Deterministic analytics

The deterministic layer performs arithmetic, aggregation, reconciliation and defined statistical analysis. It should be version-controlled and tested. Inputs and outputs retain hashes and timestamps. Material calculations have expected totals, reasonableness checks and reconciliation status.

The language model can propose code or formula logic. Approved code executes the calculation. A specialist reviews the method and material result. This boundary addresses a central risk: fluent explanations can make an incorrect number appear credible.

8.4 Generation and approval

The generation layer produces drafts from authorised evidence and deterministic outputs. Prompts specify the question, evidence perimeter, required citations, prohibited inference and output schema. Model and configuration are recorded. Unsupported statements, source conflicts and insufficient evidence should trigger exceptions.

NIST describes confabulation as plausible output that can be factually inaccurate or internally inconsistent and notes that generated citations can be false [8]. It also identifies automation bias, where people place excessive reliance on automated output. A reviewer interface should therefore show evidence before stylistic polish and require an explicit conclusion status.

8.5 Exception queue

[Matchpoint synthesis] Exceptions include missing source populations, failed tie-outs, contradictory evidence, low-confidence extraction, unmatched entities, material outliers, failed tests, expired sources, permission violations and unresolved legal or accounting questions. Each exception has severity, owner, due date, status, resolution type and evidence.

Permitted resolution types are corroborated, corrected, de-scoped, accepted, escalated or open. A report should disclose material open or accepted exceptions. Deleting an exception from the interface is not a resolution.

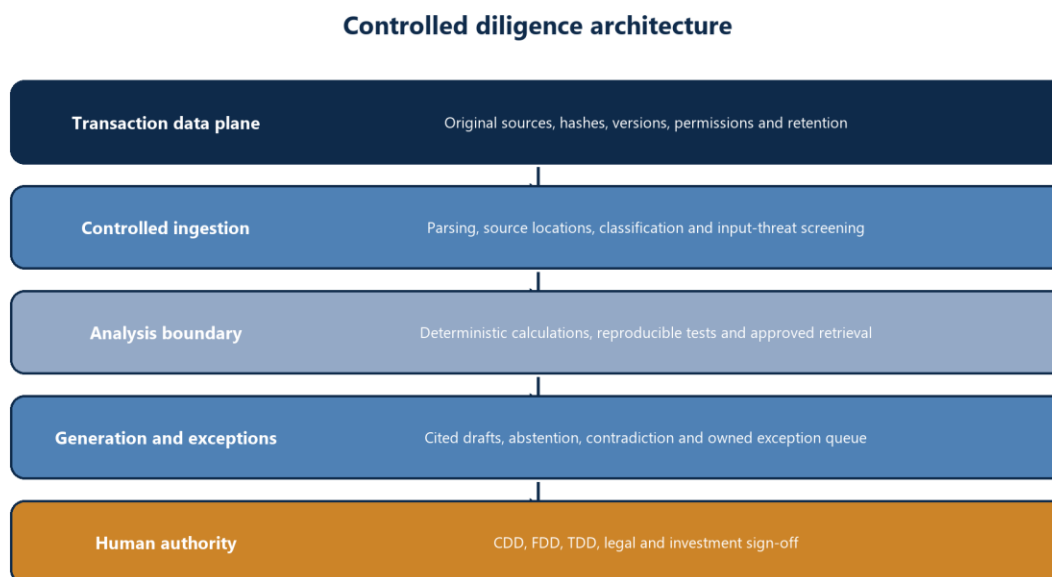


Figure 4. Controlled diligence architecture with separated data, analysis, generation, exception and human-authority layers

Source: Matchpoint architecture; AI-risk and secure-development controls informed by [7-10].

9. RISK AND CONTROL MAP

9.1 Confabulation and citation error

Generated text can contain unsupported facts, logic or citations [8]. Controls include retrieval from an approved corpus, claim-level citations, source opening, deterministic facts, abstention rules, evaluation sets and reviewer sign-off. Material conclusions should fail closed when required evidence is missing.

9.2 Scope and population error

An apparently complete result can be based on an incomplete source set. Controls include entity-period manifests, source counts, expected-population reconciliation, file-hash inventory, duplicate detection and change alerts. Coverage measures use the expected population as denominator.

9.3 Numerical error

Language models can misapply arithmetic, units or periods. Controls include deterministic calculation, type and unit checks, control totals, independent recalculation and versioned logic. Report prose cites the calculation record.

9.4 Privacy, confidentiality and privilege

Transaction rooms contain sensitive information. Controls include deal isolation, least privilege, approved endpoints, retention settings, regional and transfer assessment, privilege protocols, redaction where authorised, logs and incident response. Counsel determines legal scope.

9.5 Prompt injection and malicious content

Documents and web sources can contain instructions designed to redirect an AI system. Controls include content sanitisation, separation of data and instructions, tool allow-lists, constrained permissions, output validation and human review. A retrieval source does not receive execution authority.

9.6 Model and vendor change

Model behaviour, provider terms and components can change. Controls include inventory, version or release recording, evaluation on change, approved alternatives, concentration assessment and rollback procedures. NIST and NCSC materials support lifecycle and supply-chain governance [7-10].

9.7 Human over-reliance

NIST identifies automation bias as a human-AI configuration risk [8]. Controls include reviewer training, evidence-first interfaces, independent checks for high-risk assertions, explicit uncertainty and audit of approval behaviour. Senior reviewers should be accountable for defined conclusion families.

Risk	Failure mode	Primary control	Evidence retained
Confabulation	Unsupported fact or citation enters report	Claim-level citation and review	Prompt, output, source location, reviewer
Population gap	Missing entity, period, file or record	Expected-population reconciliation	Manifest, counts, hashes, exception
Numerical error	Incorrect arithmetic or units	Deterministic calculation and tie-out	Code or formula version, inputs, outputs, check
Privacy breach	Unauthorised disclosure or processing	Deal isolation and approved endpoint	Access log, configuration, incident record
Prompt injection	Source redirects system behaviour	Data-instruction separation and tool constraint	Scan result, retrieval record, blocked action
Model drift	Changed behaviour alters outputs	Versioned evaluation and change gate	Model version, evaluation result, approval
Automation bias	Reviewer accepts polished weak evidence	Evidence-first review and independent check	Review actions, exception decision, sign-off

10. HUMAN REVIEW AND DECISION RIGHTS

10.1 Named authority

[Matchpoint synthesis] Every material assertion should have a preparer, reviewer and conclusion owner. The preparer may operate the AI workflow. The reviewer examines evidence and method. The conclusion owner accepts the statement for the report. Roles can be combined for low-risk work under the mandate. High-risk conclusions may require an independent or senior reviewer.

10.2 Professional scepticism in the interface

The workflow should make contradiction visible. PCAOB standards describe evidence that supports and contradicts management assertions and link professional scepticism to critical assessment [4,5]. [Matchpoint synthesis] An interface should display supporting evidence, contrary evidence, unresolved gaps and source reliability together. A single generated narrative can otherwise suppress disagreement.

10.3 Decision rights by output type

Output	System authority	Human authority
Extracted fact	Proposed fact with source pointer	Reviewer validates material facts
Financial calculation	Approved deterministic engine	Financial specialist approves method and interpretation

Output	System authority	Human authority
Technical test result	Approved test tool and captured output	Technology or cyber specialist assesses meaning and materiality
Legal characterisation	Research support only	Qualified counsel
Accounting or tax treatment	Research and schedule support only	Qualified accounting or tax specialist
Investment conclusion	No autonomous authority	Investment committee or delegated decision-maker
Transaction, payment or binding communication	No autonomous authority	Authorised person under mandate

10.4 Sign-off certificate

[Matchpoint synthesis] The report release record should identify the mandate version, source cut-off, material open exceptions, model and tool versions, calculation and test versions, reviewers, approval time and distribution. If the data room changes after cut-off, forward traceability identifies conclusions that need refresh.

11. PILOT DESIGN AND MEASUREMENT

11.1 Baseline before automation

A pilot requires a baseline that uses the same work definition. The team should select a bounded module, define its expected source population, procedures, outputs and materiality, and measure current elapsed and reviewer time. Historical time records can be incomplete or inconsistent. Where the baseline relies on management estimates, it should be labelled '[Unverified management estimate]' until observed.

11.2 Shadow-mode pilot

[Matchpoint synthesis] Shadow mode runs the AI-assisted workflow alongside the authorised conventional process. The assisted output does not control the live report. Reviewers compare source coverage, facts, calculations, exceptions and conclusions. A labelled benchmark set is needed for accuracy and false-negative estimates.

Suitable first modules include contract-field extraction, customer-name normalisation, data-room manifesting, recurring adjustment evidence indexing or dependency inventory. The module should be material enough to test controls and bounded enough to review completely.

11.3 Measurement dictionary

Measure	Definition	Minimum record	Interpretation
Source coverage	In-scope sources processed / expected sources	Manifest and exclusions	Coverage does not equal correctness
Citation validity	Sampled material claims with correct source location / sampled claims	Claim sample and review result	Report by materiality tier
Calculation fidelity	Reperformed material outputs matching approved result / tested outputs	Test set and tolerance	Deterministic layer target should be exact within defined tolerance
Exception precision	Confirmed material exceptions / reviewed exceptions	Labelled review set	Low precision consumes reviewer time
Exception recall	Confirmed benchmark exceptions surfaced / total benchmark exceptions	Labelled benchmark	Unknown issues outside benchmark remain possible

Measure	Definition	Minimum record	Interpretation
Reviewer time	Human minutes for setup, review, correction and approval	Time record by assertion	Separate from machine time
Elapsed time	Completion timestamp minus defined start	Start and finish events	Identify external waiting time
Rework	Reviewer minutes spent correcting assisted output	Correction log	Classify extraction, analysis and drafting rework
Issue yield	Decision-useful issues accepted by specialists	Issue register	Requires a defined materiality rule

11.4 Stage gates

[Matchpoint synthesis] A pilot advances only when approved gates are met. Gate 1 confirms data and legal authority. Gate 2 confirms source and calculation controls. Gate 3 confirms accuracy and exception performance on the benchmark. Gate 4 confirms reviewer workflow and incident readiness. Gate 5 authorises controlled production for a named module. Gate 6 reviews production evidence before scope expansion.

The pilot should define stop conditions such as unauthorised disclosure, material calculation divergence, inaccessible evidence, unstable model behaviour or unacceptable false-negative performance on the benchmark. Stop conditions are operational controls and require local approval.

Diligence risk and control map

DOMAIN	EXPOSURE	CONTROL OBJECTIVE
Evidence	Population gap, stale source	Manifest, version and scope reconciliation
Model	Confabulation, omission, drift	Citations, evaluation, abstention and change gate
Analysis	Arithmetic or test error	Deterministic calculation and reproducible test
Security	Leakage, injection, access	Isolation, least privilege, monitoring and incident record
Human	Over-reliance, weak review	Evidence-first interface and named sign-off

Figure 5. Risk and control map across evidence, model, analysis, security and human review

Source: Matchpoint synthesis based on [4-10].

12. BUY, BUILD OR OUTSOURCE

12.1 Decision dimensions

The operating choice depends on confidentiality, source complexity, integration, custom methods, internal engineering capacity, assurance needs, vendor concentration, speed and total cost. A commercial data-room or diligence platform may offer mature ingestion and collaboration. Internal development may support proprietary methods and integration. A service provider may combine technology with specialist review.

Provider material can describe capabilities. It does not independently establish accuracy or economic value. PwC's technology-diligence service page is one practitioner description of current domains [18]. A buyer should evaluate the actual product, deployment, controls, contract and evidence against its mandate.

Dimension	Buy a platform	Build internally	Outsource a managed workflow
Time to bounded pilot	Often shorter after procurement and configuration	Depends on data and engineering readiness	Depends on provider onboarding and data access
Method customisation	Product and configuration limits	Highest potential flexibility	Negotiated within service method
Data control	Depends on architecture and contract	Direct internal control subject to capability	Shared operating and contractual control
Specialist capacity	Usually separate	Must be staffed internally	Can be integrated into service
Model and vendor concentration	Platform dependency	Component dependencies remain	Provider and platform dependencies
Evidence ownership	Export and audit rights require diligence	Can be designed into system	Must be explicit in contract and delivery

12.2 Vendor diligence questions

Vendor diligence should cover data retention and training use; tenancy and isolation; encryption and keys; access administration; subprocessors; regions and transfers; model and component inventory; prompt and output logging; version changes; evaluation; incident history; resilience; deletion; export; audit rights; intellectual property; liability; and termination support. AI vendors that develop models should also be assessed against relevant secure-development and supply-chain practices [9,10].

12.3 Economic case

[Matchpoint synthesis] The economic case should compare implementation and operating cost with reviewer time, elapsed time where economically relevant, coverage expansion, avoided rework and decision-useful issue yield. Licence cost alone understates the operating requirement. Control design, data preparation, specialist review, evaluation and change management are recurring costs.

13. MATCHPOINT IMPLEMENTATION ROADMAP

13.1 Phase 0: mandate and authority

Define the transaction module, jurisdictions, data classes, reliance, materiality, output users and decision rights. Obtain information-security, privacy, legal and engagement approvals. Record prohibited systems and prohibited actions.

13.2 Phase 1: assertion and evidence foundation

Create the assertion map, source manifest, evidence taxonomy, exception categories and sign-off roles. Establish deterministic calculations or reproducible tests for the selected module. Measure the conventional baseline.

13.3 Phase 2: shadow pilot

Run the assisted workflow without production authority. Build a labelled evaluation set from representative sources and known edge cases. Review every material output. Record citations, errors, missed exceptions, reviewer time and rework.

13.4 Phase 3: controlled production

Authorise a named module, user group, data boundary and model version. Keep report conclusions under human sign-off. Monitor quality and security measures. Preserve the conventional fallback.

13.5 Phase 4: cross-workstream integration

Connect CDD, FDD and TDD assertions where the investment thesis depends on shared facts. Examples include customer retention and revenue quality; product roadmap and forecast growth; cloud architecture and gross margin; cyber remediation and net debt or price adjustment. Maintain specialist conclusion ownership.

13.6 Phase 5: portfolio learning

Store de-identified or authorised precedent structures, issue taxonomies and evaluation cases subject to confidentiality and contract. Compare process performance across engagements using consistent definitions. Portfolio learning should not expose one transaction's confidential evidence to another.

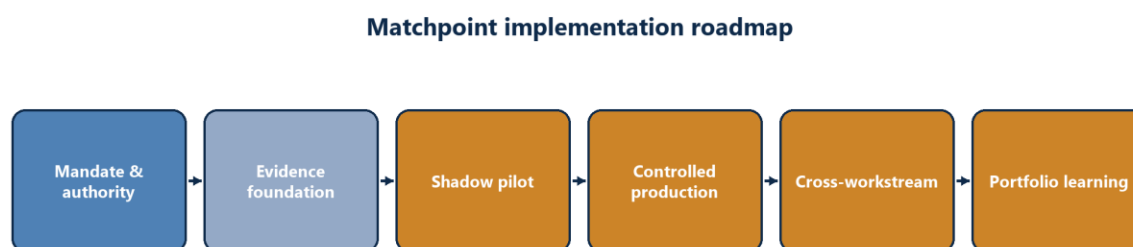


Figure 6. Matchpoint implementation roadmap from mandate and baseline to controlled production and portfolio learning

Source: Matchpoint scenario analysis. Timing, resources and stage-gate thresholds require separate approval.

14. MATCHPOINT SCENARIO ANALYSIS

14.1 Scenario status

The following scenario is '[Matchpoint scenario analysis]'. It is illustrative, unverified and does not describe an observed Matchpoint engagement. Values require replacement with an approved baseline and pilot results before any commercial claim.

14.2 Illustrative module

Assume a mid-market acquisition workstream receives 1,200 source objects across CDD, FDD and TDD. The pilot selects three bounded modules: customer-contract field extraction; EBITDA-adjustment evidence indexing; and software-dependency inventory. The conventional work programme and materiality remain fixed. Specialists retain every conclusion.

The assisted workflow registers all received objects, performs controlled extraction, links proposed facts to source locations and routes exceptions. Deterministic code calculates customer concentration and adjustment bridges. Approved tools produce the dependency inventory. Reviewers inspect all high-materiality outputs and a defined sample of lower-materiality outputs.

14.3 Illustrative scorecard

Measure	Baseline	Pilot result	Status
Expected source population	'[Unverified management estimate: 1,200 objects]'	To be measured	No result claimed

Measure	Baseline	Pilot result	Status
Source coverage	To be measured	To be measured	Gate requires reconciled manifest
Citation validity	To be measured	To be measured	Gate set by approved materiality tier
Calculation fidelity	To be measured	To be measured	Material results require deterministic match
Reviewer hours	To be measured	To be measured	Include setup and rework
Elapsed days	To be measured	To be measured	Record external waiting separately
Confirmed material issues	To be measured	To be measured	Quality and relevance review required
Material open exceptions	To be measured	To be measured	Disclose at release

14.4 Decision rule

[Matchpoint synthesis] Production approval requires satisfactory evidence quality, information-security approval, acceptable reviewer workload and tested incident response. A time saving does not override a failed quality or security gate. Coverage expansion can justify adoption even where elapsed time remains unchanged, subject to an approved economic case.

15. LIMITATIONS AND RESEARCH AGENDA

This paper has six material limitations. First, the most directly relevant adoption and efficiency evidence is self-reported survey evidence. It does not independently observe transaction cycle time, error, quality or returns. Second, the Thomson Reuters sample spans several professions and does not isolate the headline metrics to transaction diligence. Third, the paper uses PCAOB audit principles by analogy; an FDD mandate is not a statutory audit and may have different objectives, evidence thresholds and reliance.

Fourth, AI systems, provider terms and enterprise controls change. A paper dated July 2026 cannot substitute for current vendor and configuration diligence. Fifth, legal, privacy, accounting, tax, competition and disclosure requirements depend on parties, transaction structure and jurisdiction. This paper is not professional advice. Sixth, the Matchpoint operating model and scenario are analytical proposals. Their productivity and quality effects require controlled pilot evidence.

Further research should report quality-adjusted time by diligence module; compare population-wide extraction with specialist sampling; construct labelled exception benchmarks; study reviewer over-reliance and interface design; measure evidence-graph maintenance cost; analyse cross-border data-room architectures; and examine whether earlier issue visibility changes price, structure, conditions or post-close plans. Transaction-return analysis would require careful controls for selection, market and execution factors.

16. CONCLUSION

AI-assisted diligence is already visible in reported market practice. KPMG's survey indicates adoption in diligence and valuation and respondent-reported efficiency in analytical tasks [1,2]. Thomson Reuters documents broad professional-services adoption and a much smaller reported base of formal ROI measurement [3]. The combination supports a disciplined implementation question: how should the workflow be designed and measured?

[Matchpoint synthesis] The answer begins with assertions and evidence. Register the source population. Preserve provenance and contradiction. Use AI for controlled ingestion, extraction, comparison, triage and drafting. Use deterministic systems for financial calculations and reproducible tools for technical tests. Route uncertainty and failures into an owned exception queue. Assign specialists and authorised decision-makers to conclusions.

The resulting system can be tested. Source coverage, citation validity, calculation fidelity, exception performance, reviewer time, elapsed time, rework and issue yield create a measurable record. Claims about compression or quality can then be grounded in the named pilot. Until that evidence exists, they remain hypotheses.

DECLARATIONS

Evidence classification

Survey observations are reported with sample and method boundaries. PCAOB standards are used as an audit-evidence analogy for FDD and retain their audit-scope caveat. NIST, NCSC, DOJ, FATF, IFRS, SEC and OECD materials are applied within the stated boundaries. Matchpoint frameworks and scenarios are labelled.

Conflicts and funding

No external funding was received for this paper. Matchpoint Partners may provide advisory services related to technology, transactions, diligence, governance and implementation. References to provider materials do not constitute endorsement.

Advice and reliance

This paper is research only. It is not legal, tax, accounting, cyber-security or investment advice. Readers should obtain qualified advice for the relevant transaction and jurisdiction.

REFERENCES

- [1] KPMG LLP, *2026 Global M&A Outlook: The Year of the Carve-Out*, March 2026. Survey of 700 senior M&A decision-makers across 20 countries and jurisdictions, fielded 19 December 2025 to 27 January 2026. <https://kpmg.com/kpmg-us/content/dam/kpmg/pdf/2026/kpmg-2026-global-ma-outlook-final.pdf>
- [2] KPMG International, "KPMG International survey of 700 global dealmakers reveals rising M&A expectations," 18 March 2026. <https://kpmg.com/xx/en/media/press-releases/2026/03/kpmg-survey-of-global-dealmakers-reveals-rising-m-and-a-expectations.html>
- [3] Thomson Reuters Institute, *2026 AI in Professional Services Report*, 2026. Online survey of 1,514 respondents in 27 countries, conducted October and November 2025. <https://www.thomsonreuters.com/content/dam/ewp-m/documents/thomsonreuters/en/pdf/reports/2026-ai-in-professional-services-report.pdf>
- [4] Public Company Accounting Oversight Board, AS 1105: Audit Evidence. <https://pcaobus.org/oversight/standards/auditing-standards/details/AS1105>
- [5] Public Company Accounting Oversight Board, AS 2301: The Auditor's Responses to the Risks of Material Misstatement. <https://pcaobus.org/oversight/standards/auditing-standards/details/AS2301>
- [6] Public Company Accounting Oversight Board, Release No. 2023-004, *Proposed Amendments Related to Aspects of Designing and Performing Audit Procedures that Involve Technology-Assisted Analysis of Information in Electronic Form*, 26 June 2023; and PCAOB technology-assisted analysis standards update, 2024. <https://assets.pcaobus.org/pcaob-dev/docs/default-source/rulemaking/docket-052/pcaob-release-no.-2023-004-technology-assisted-analysis.pdf>
- [7] E. Tabassi, *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*, NIST AI 100-1, National Institute of Standards and Technology, 2023. <https://doi.org/10.6028/NIST.AI.100-1>
- [8] C. Autio et al., *Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile*, NIST AI 600-1, National Institute of Standards and Technology, 2024. <https://doi.org/10.6028/NIST.AI.600-1>
- [9] H. Booth et al., *Secure Software Development Practices for Generative AI and Dual-Use Foundation Models: An SSDF Community Profile*, NIST SP 800-218A, 2024. <https://doi.org/10.6028/NIST.SP.800-218A>
- [10] UK National Cyber Security Centre et al., *Guidelines for Secure AI System Development*, November 2023. <https://www.ncsc.gov.uk/files/Guidelines-for-secure-AI-system-development.pdf>
- [11] US Department of Justice and Federal Trade Commission, *2023 Merger Guidelines*, 18 December 2023. <https://www.justice.gov/atr/merger-guidelines>
- [12] Financial Action Task Force, *Guidance on Beneficial Ownership of Legal Persons*, 10 March 2023. <https://www.fatf-gafi.org/en/publications/Fatfrecommendations/Guidance-Beneficial-Ownership-Legal-Persons.html>
- [13] UAE Legislation, Federal Decree-Law No. 45 of 2021 Concerning the Protection of Personal Data. <https://uaelegislation.gov.ae/en/legislations/1972>
- [14] IFRS Foundation, IFRS 3 Business Combinations. <https://www.ifrs.org/issued-standards/list-of-standards/ifrs-3-business-combinations/>
- [15] US Securities and Exchange Commission, *Cybersecurity Risk Management, Strategy, Governance, and Incident Disclosure: A Small Entity Compliance Guide*. <https://www.sec.gov/resources-small-businesses/small-business-compliance-guides/cybersecurity-risk-management-strategy-governance-incident-disclosure>
- [16] OECD, *OECD Due Diligence Guidance for Responsible Business Conduct*, OECD Publishing, Paris, 2018. <https://doi.org/10.1787/15f5f4b3-en>
- [17] PwC, *M&A Outlook 2026: Expectations Are High for the Year Ahead*, 2026. <https://www.pwc.com/gx/en/1/issues/c-suite-insights/the-leadership-agenda/ma-outlook.html>
- [18] PwC US, *AI and Technology Due Diligence*. <https://www.pwc.com/us/en/services/consulting/deals/strategy-value-creation/ai-technology-diligence.html>

APPENDIX A. ASSERTION CONTROL RECORD

A.1 Minimum fields

Field	Purpose
Assertion ID and wording	Atomic statement being tested
Workstream and owner	CDD, FDD, TDD or cross-workstream accountability
Materiality and risk	Determines evidence and reviewer threshold
Scope and perimeter	Entity, period, geography, product, customer or environment
Expected evidence	Required source classes and procedures
Evidence received	Registered source and version identifiers
Contrary evidence	Sources or results that challenge the assertion
Calculation or test	Versioned deterministic or reproducible procedure
Exceptions	Open, accepted, escalated, corrected, corroborated or de-scoped
Conclusion and limitation	Approved wording and known boundary
Reviewer and timestamp	Named authority and decision time

A.2 Conclusion status

- **Supported:** the approved evidence threshold is met.
- **Partially supported:** some elements meet the threshold and stated limitations remain.
- **Unsupported:** available evidence does not meet the threshold.
- **Contradicted:** relevant evidence conflicts materially with the assertion.
- **Open:** required work or evidence remains outstanding.
- **Out of scope:** the mandate does not authorise a conclusion.

APPENDIX B. EXCEPTION RECORD

Field	Description
Exception ID	Stable identifier
Source or procedure	Originating object, calculation or test
Assertion impact	Conclusions potentially affected
Type	Missing, contradictory, failed tie-out, low confidence, access, legal or other
Severity	Mandate-specific materiality classification
Owner and due date	Responsible person and target resolution time
Resolution	Corroborated, corrected, de-scoped, accepted, escalated or open
Evidence and approval	Supporting record and authorised reviewer

APPENDIX C. MODEL, PROMPT AND TOOL RECORD

Field	Description
System and version	Model, retrieval system, parser, code or test tool
Provider and deployment	Contracting entity, endpoint, region and tenancy
Configuration	Temperature, tool permissions, retrieval scope and policy
Prompt or procedure version	Reproducible instruction or test identifier
Input population	Source manifest and access class
Evaluation result	Benchmark date, sample and result
Output and exceptions	Stored result, citations and routed issues
Reviewer	Named approval or rejection

APPENDIX D. RELEASE CERTIFICATE

Field	Release evidence
Mandate and report version	Approved scope and deliverable identifier
Source cut-off	Latest included source timestamp
Material open exceptions	Listed or explicitly none within the recorded scope
Model and tool versions	Production versions used
Calculation and test versions	Approved analytical artefacts
Workstream approvals	CDD, FDD and TDD reviewers as applicable
Legal and specialist qualifications	Scope limitations and referrals
Distribution and timestamp	Authorised recipients and release time

APPENDIX E. GLOSSARY

Assertion. An atomic statement tested through a defined evidence and review procedure.

CDD. Commercial due diligence as defined for this paper; distinct from regulatory customer due diligence.

Compression. Reduced elapsed or reviewer time against a defined baseline, or explicitly labelled coverage expansion within comparable resources.

Evidence graph. Structured links among assertions, sources, versions, facts, transformations, calculations, tests, exceptions, conclusions and reviewers.

FDD. Financial due diligence; it is not a statutory audit.

Matchpoint scenario analysis. An illustrative, unverified analytical case that does not describe observed client results.

Rigour. Satisfaction of the evidence, calculation, testing, review and exception rules approved for the mandate.

TDD. Technology due diligence, including AI-system diligence where relevant.