

THE FOUR-PAGE FRAMEWORK FOR FINANCIAL SERVICES FIRMS

SECURE USE OF AI

You can outsource the tool. **You cannot outsource the obligation.** Every AI your firm uses, buys, or gets attacked by still lands on your supervision, your books and records, and your client protection duties.

1 · UNDERSTAND

2 · DECIDE

3 · COMPLY

4 · DO

ALIGNED TO NIST AI RMF · OWASP LLM TOP 10 · SANS · ISO 42001

1 in 4

malicious breaches were AI-enabled in 2026 — a 56% jump in one year

\$6.29M

average breach cost in financial services — second-highest of any industry

43%

of incidents involved shadow AI — unapproved tools staff used anyway. It was 20% last year

92%

of firms that suffered an AI-related breach had no real access controls on AI

45%

of AI-driven attacks were deepfake or impersonation — voice, video, and email

01 THE THREE FACES OF AI RISK

CONCEPTUAL

Most firms write one AI policy and think they are done. AI arrives in your firm through three separate doors, and each door needs a different control. Miss one and the policy is decorative.

DOOR 01

AI YOUR PEOPLE USE

Copilot, ChatGPT, Claude, meeting note-takers, and whatever an advisor installed last Tuesday.

WHAT ACTUALLY GOES WRONG

- ◆ Client PII, account numbers, and statements pasted into a public tool
- ◆ A confident, wrong answer sent to a client without anyone checking it
- ◆ AI-drafted marketing that becomes an unsubstantiated performance claim
- ◆ Business records created in a tool your firm cannot search or retain

WHO OWNS IT

CCO and the business line manager — not IT alone.

First control: an approved tool list with enterprise terms, plus a hard rule that client data never touches an unapproved tool.

DOOR 02

AI YOUR VENDORS USE

Your custodian, CRM, planning software, portfolio tools, email security, and your MSP all shipped AI features you never bought.

WHAT ACTUALLY GOES WRONG

- ◆ Your client data trains someone else's model because the default said yes
- ◆ A vendor sub-processes to a model provider you have never assessed
- ◆ An AI feature turns on by default and nobody in compliance is told
- ◆ The vendor breaches — and it is still your Reg S-P notification

WHO OWNS IT

Vendor management, with security review support.

First control: add AI questions to your vendor due diligence — training rights, data residency, sub-processors, and opt-out.

DOOR 03

AI USED AGAINST YOU

Attackers got the same tools you did, and they are using them on your staff and your clients right now.

WHAT ACTUALLY GOES WRONG

- ◆ A cloned voice — built from three seconds of audio — authorizes a wire
- ◆ Deepfake video on a call that looks and sounds exactly like your principal
- ◆ Flawless phishing with no typos, in your CEO's actual writing style
- ◆ Prompt injection: hidden instructions in a document that hijack your AI tool

WHO OWNS IT

Operations and the money-movement process owner.

First control: out-of-band verification on every funds and banking change. Voice and video are no longer identity.

02 FIVE QUESTIONS AN EXAMINER WILL ASK YOU

EXAM READINESS

If you can answer these five with a document rather than a sentence, you are in good shape. If you answer with a sentence, you have a finding.

Q1 · INVENTORY

SHOW ME EVERY AI TOOL IN USE HERE

Including the ones embedded in software you already licensed, and the ones staff signed up for on their own.

Q2 · AUTHORITY

WHO APPROVED THIS USE CASE, AND ON WHAT BASIS?

A named approver, a risk tier, and a date. Not "we discussed it."

Q3 · DATA

WHAT CLIENT DATA CAN REACH THE MODEL?

And what stops it — technically, not just in policy language.

Q4 · SUPERVISION

WHO REVIEWS AI OUTPUT BEFORE A CLIENT SEES IT?

Your supervisory procedures have to name the reviewer and the trigger.

Q5 · RECORDS

WHERE ARE THE AI-GENERATED RECORDS RETAINED?

Books-and-records rules did not change because the author was a model.

THE OTHER SIDE OF THE LEDGER

The same technology cuts both ways. Firms using AI and automation extensively in their security operations cut roughly **\$1.93M off the cost of a breach** and contained incidents **65 days faster** than firms that did not. Governed AI is not just a compliance obligation — it is one of the few controls that pays for itself. The goal of this framework is not to slow your firm down. It is to let your people use AI at full speed, defensibly.

03 THE ASCENT: SIX STAGES FROM BASE CAMP TO AUDIT-READY

STRATEGIC · 90–120 DAYS

Do them in order. Each stage produces a document — that document is what you hand an examiner. If a stage produces nothing you can print, the stage is not finished.

STAGE 01	STAGE 02	STAGE 03	STAGE 04	STAGE 05	STAGE 06
GOVERN Who decides what AI this firm may use, and who is accountable when it goes wrong? OWNER CCO with executive sponsor PRODUCES AI acceptable use policy, named AI owner, approval workflow Done when: a person's name — not a committee — signs AI approvals.	INVENTORY What AI is actually running here — including inside tools you already pay for? OWNER IT / MSP with compliance PRODUCES AI asset register: tool, owner, data touched, risk tier, approval date Done when: the register includes vendor-embedded AI, not just chatbots.	PROTECT What technically prevents client data from reaching an unapproved model? OWNER IT / security PRODUCES DLP rules, SSO and MFA on AI tools, browser or CASB controls, tenant config Done when: a test upload of fake client data is actually blocked.	SUPERVISE Who reviews AI output before it reaches a client, a regulator, or a trade? OWNER Supervisory principal PRODUCES Updated WSPs naming the reviewer, the trigger, and the evidence of review Done when: your WSPs mention AI by name and describe the review step.	MONITOR How would you know if someone put a client list into a public AI tool last Thursday? OWNER IT / MSP with compliance PRODUCES AI usage logging, shadow AI discovery, exception reporting, quarterly review Done when: you can produce a log, not an opinion.	RESPOND What happens in the first hour after AI misuse, a deepfake wire, or a model leak? OWNER Incident response lead PRODUCES IR plan with AI scenarios, Reg S-P notification path, tested tabletop Done when: you have run the scenario, not just written it.

04 TIER THE USE CASE, NOT THE TOOL

APPROVAL MODEL

The same AI tool can be low risk for one task and unacceptable for the next. Approve use cases. This table is the fastest way to give your people a yes without giving away the firm.

TIER	WHAT IT LOOKS LIKE IN A WEALTH OR BROKERAGE FIRM	APPROVAL PATH	WHAT IS REQUIRED BEFORE ANYONE USES IT
CRITICAL	AI that recommends or executes: trade suggestions, suitability screening, automated client responses, anything that moves money or makes an agentic decision on its own.	Executive + CCO sign-off, documented	Written risk assessment, human approval on every output, full audit trail, vendor model documentation, and a tested rollback. If you cannot explain the output to a client, do not deploy it.
HIGH	AI touching client PII or account data: CRM summarizers, meeting note-takers on client calls, document analysis of statements, planning tools with AI features.	CCO approval, annual review	Enterprise agreement with no-training terms, DLP and SSO enforced, retention mapped to books-and-records, named supervisory reviewer, client-consent language reviewed.
MODERATE	Internal work product using firm-confidential but non-client data: policy drafts, RFP responses, internal analysis, training content, code.	Manager approval from the approved-tool list	Approved tool only, no client identifiers, human review before anything is finalized or distributed, and the output is retained where your firm can find it.
LOW	Public-information work with no firm or client data: general research, brainstorming, summarizing published articles, drafting a job description.	Standing approval — just train people	Approved tool, annual training, and the same rule as everything else: a human is accountable for what goes out the door.

05 WHAT WE FIND IN ALMOST EVERY ASSESSMENT

FIELD NOTES

These eight gaps show up again and again at firms that already believe they have AI covered. Read them as a pre-check before you spend money on anything else.

- The policy exists; the inventory does not.** A firm can produce a signed AI policy but cannot list which tools are actually in use. If an AI asset is missing from your inventory, it is missing from your governance.
- Vendor-embedded AI was never assessed.** The CRM, the planning tool, and the email gateway all added AI features. None went through vendor due diligence, because nobody bought anything new.
- Personal accounts doing firm work.** Staff use free consumer tiers on personal logins, so there is no enterprise agreement, no SSO, no logging, and no way to prove what left the building.
- Supervisory procedures never mention AI.** The WSPs describe email and social media review in detail and are silent on the tool drafting half the firm's client communications.
- Money movement still trusts a voice.** Callback procedures exist but allow the number from the request, or let a senior person waive the second approver. Both are the exact opening a deepfake needs.
- No logging, so no answer.** Asked whether client data ever went into a public AI tool, the honest answer is "we don't think so." That is not an answer an examiner or an insurer accepts.
- Incident response has no AI scenario.** The plan covers ransomware and email compromise. It says nothing about a model leak, a prompt injection, or a deepfake-authorized transfer.
- Training is annual, generic, and untested.** One slide on AI in a compliance deck. No role-based content for the people who actually approve wires, and no tabletop to see whether it stuck.

06 WHO IS WATCHING, AND WHAT THEY EXPECT

REGULATORY MAP

No US regulator has published a standalone AI rulebook for advisers and broker-dealers. That is the trap: your existing obligations already cover AI, which means you are already accountable.

SEC

RIAS & BROKER-DEALERS

- **Reg S-P amendments** — written incident response program and customer notification within 30 days, including when a vendor is the one breached
- **Marketing Rule** — AI claims about your services must be substantiated. "AI-washing" is an enforcement theme
- **Books and records** — AI-generated communications are still records
- **Exam priorities** — expect questions on AI use, governance, and third-party oversight

FINRA

MEMBER BROKER-DEALERS

- **Rule 3110 supervision** — your WSPs must cover the technology your people actually use
- **Notice 24-09** — GenAI is covered by existing rules; supervise it, test it, and govern the data
- **Annual Oversight Report** — AI, cybersecurity, and third-party risk are recurring focus areas
- **Rule 2210** — AI-drafted communications with the public still need principal review

NYDFS

PART 500 COVERED ENTITIES

- **AI cybersecurity guidance** — AI-enabled social engineering and deepfakes must be addressed in your risk assessment
- **500.9** — risk assessments must be updated when AI changes your risk profile
- **500.11** — third-party policy must cover AI vendors and their sub-processors
- **500.12 / 500.14** — MFA and training that specifically covers AI-driven fraud

FRAMEWORKS

USE THESE TO SHOW YOUR WORK

- **NIST AI RMF 1.0 + GenAI Profile** — the govern / map / measure / manage structure examiners recognize
- **OWASP Top 10 for LLM Apps** — the technical risk list: prompt injection, data disclosure, supply chain
- **SANS Critical AI Security Guidelines** — practical control baseline
- **ISO/IEC 42001** — if you want a certifiable AI management system

07 TWELVE CONTROLS, AND THE EVIDENCE EACH ONE PRODUCES

TACTICAL CHECKLIST

Tick the box only if you could hand over the evidence today. Count your ticks and use the scale at the bottom — this is the same scoring we apply on a client assessment.

01 ☐ **Published approved-AI-tool list**

One list, one named owner, and a route for staff to request additions so they stop going around you.

[Approved AI Tools Register](#)

03 ☐ **AI asset inventory, refreshed quarterly**

Every model, tool, agent, and AI feature — including the ones bundled into software you already own.

[AI Asset Inventory](#)

05 ☐ **SSO and MFA on every AI tool**

Phishing-resistant MFA for anyone with admin rights or the ability to move money. This is the control 92% of breached firms were missing.

[Identity provider config export](#)

07 ☐ **Shadow AI discovery report**

A monthly view of unsanctioned AI use from your browser controls, CASB, or firewall. Assume it exists until the report says otherwise.

[Monthly shadow AI report](#)

09 ☐ **Supervisory procedures that name AI**

Your WSPs should state which AI outputs get reviewed, by whom, when, and how that review is recorded.

[WSP redline + review log](#)

11 ☐ **Out-of-band verification for money movement**

A written, mandatory callback and dual-approval procedure for every wire, banking change, and address change. The drill is on page 4.

[Written procedure + verification log](#)

02 ☐ **AI acceptable use policy, attested annually**

Plain-language rules on what may and may not be entered into an AI tool, signed by every employee and contractor.

[Policy + signed attestations](#)

04 ☐ **Business-tier terms with no-training rights**

Free and personal accounts are not a control environment. Contract for it: no training on your data, defined retention, named sub-processors.

[Executed agreement + DPA](#)

06 ☐ **DLP that blocks client data to AI destinations**

Policy tells people what not to do. DLP stops them. Test it with fabricated client data and keep the result.

[DLP rule set + block test result](#)

08 ☐ **AI questions in vendor due diligence**

Training rights, data location, sub-processors, opt-out, breach notification, and whether AI features default to on.

[Updated VRA questionnaire](#)

10 ☐ **Retention mapped for AI-generated records**

If AI helped draft a client communication or a business record, it is captured, searchable, and retained like any other.

[Retention matrix + archive test](#)

12 ☐ **Role-based training and a tested tabletop**

Advisors, ops, and finance face different attacks. Run a deepfake and AI-misuse scenario at least annually and record the lessons.

[Training roster + after-action report](#)

0-4 TICKS · EXPOSED

You have adoption without governance. This is where the \$6.3M average financial services breach starts. Treat the first 30 days on this page as urgent.

5-9 TICKS · PARTIAL

You have policy but thin evidence. An examiner will find the gap between what your documents say and what your systems actually do.

10-12 TICKS · DEFENSIBLE

You can show your work. Keep the inventory current, run the tabletop, and re-score every quarter as tools and staff change.

THE EXAM FOLDER

If you assemble one evidence set this year, assemble this one. Every item above rolls up into it, and it is what we build with clients during an assessment.

[01 AI acceptable use policy + signed attestations](#)[02 AI asset inventory with owners and risk tiers](#)[03 AI risk assessment mapped to NIST AI RMF](#)[04 Use-case approval records](#)[05 Vendor due diligence files with AI addendum](#)[06 WSP excerpt naming AI review triggers](#)[07 DLP configuration + block test evidence](#)[08 Identity, SSO, and MFA configuration export](#)[09 Shadow AI monitoring reports](#)[10 Incident response plan with AI scenarios](#)[11 Tabletop after-action report](#)[12 Training roster and role-based curriculum](#)

08 DEEPFAKE AND AI FRAUD: THE VERIFY DRILL

POST THIS AT THE DESK

A voice can be cloned from roughly three seconds of audio, and 45% of AI-driven attacks are impersonation. The defense is procedural, not technical. Print this and put it where wires get approved.

VERIFY BEFORE YOU MOVE

EVERY WIRE • EVERY BANKING CHANGE • EVERY TIME • NO EXCEPTIONS FOR EXECUTIVES

- 1 Stop the clock.**
Urgency is the attack. No legitimate request is destroyed by a ten-minute pause — say so out loud to the caller.
- 2 Hang up and call back on your number.**
Use the contact in your system of record, never a number supplied in the request. Spoofed callbacks route straight to the attacker.
- 3 Ask for the code word.**
Agree a verbal passphrase with clients and executives in advance, and rotate it annually. A cloned voice does not know it.
- 4 Two people, two channels.**
A second approver confirms through a different medium than the one the request arrived on.
- 5 Verify the destination, not the person.**
Attackers can fake any face or voice. They cannot fake the receiving account — check it against your records before release.
- 6 If money already moved:**
call the bank's fraud line inside the first hour, file at ic3.gov, preserve everything, and start your Reg S-P notification assessment.

NEVER, NO MATTER WHO ASKS

- ✗ Never treat a voice or a face on a video call as proof of identity
- ✗ Never call back a number supplied inside the request itself
- ✗ Never let one person complete a funds transfer alone
- ✗ Never approve a banking-detail change on email alone
- ✗ Never assume a familiar writing style means a familiar sender
- ✗ Never let seniority override the procedure — that is the whole play

RED FLAGS IN THE MOMENT

- ! New or changed payment instructions on a relationship you have had for years
- ! The request arrives outside normal hours, or while the "sender" is known to be travelling
- ! Pressure to skip the usual approver, or secrecy framed as confidentiality
- ! The caller resists a callback, or says there is no time for one
- ! A near-miss in the sending domain, or a reply-to that does not match
- ! Audio with flat emotion, odd cadence, or a beat of latency before each answer

09 YOUR FIRST 90 DAYS

If you do nothing else on this sheet, do these in this order.

ROADMAP

30 DAYS · STOP THE BLEEDING

- ◆ Publish the approved tool list and one interim rule: no client data in unapproved AI
- ◆ Turn on SSO and MFA for every AI tool in use
- ◆ Issue the verify drill for money movement and brief every person who touches a wire
- ◆ Tell staff how to request a tool, so shadow AI has somewhere legitimate to go

60 DAYS · BUILD THE FRAME

- ◆ Complete the AI asset inventory, including vendor-embedded AI features
- ◆ Adopt the AI acceptable use policy and the four-tier approval model
- ◆ Add AI questions to vendor due diligence and to annual vendor re-reviews
- ◆ Enable DLP for client data, then test it and keep the result

90 DAYS · PROVE IT

- ◆ Run an AI risk assessment mapped to NIST AI RMF and your regulator's expectations
- ◆ Update WSPs and incident response to name AI scenarios explicitly
- ◆ Run a deepfake tabletop and document the after-action findings
- ◆ Deliver a board or CCO reporting pack with owners, dates, and residual risk

FIND OUT WHERE YOU STAND IN TWO TO THREE WEEKS

Security Basecamp has focused exclusively on cybersecurity for regulated industries since 2013 — no general IT practice and no product-vendor bias. Hundreds of risk assessments, penetration tests, and vendor risk assessments for broker-dealers, RIAs, insurance companies, private equity firms, and FinTech providers. Our mission is to be the most capable cybersecurity firm serving regulated industries, combining deep technical expertise, regulatory fluency, and practical execution.

WHAT WE DO

- ◆ Cybersecurity Risk Assessments
- ◆ Third-Party & Vendor Risk Assessments
- ◆ Vulnerability Mgmt & Penetration Testing
- ◆ Breach Response & Forensic Analysis
- ◆ vCISO / CISO Support Services
- ◆ Cloud Security — AWS, Azure, M365
- ◆ Threat Intelligence & Monitoring Advisory
- ◆ Secure Use of AI

CISSP

SECURITY+

PENTEST+

GIAC

OSCP

REGULATED INDUSTRIES ONLY

BE SECURE. BE COMPLIANT.

SOURCES: IBM / PONEMON COST OF A DATA BREACH REPORT 2026 • NIST AI RMF 1.0 AND GENERATIVE AI PROFILE • OWASP TOP 10 FOR LLM APPLICATIONS • SANS CRITICAL AI SECURITY GUIDELINES • ISO/IEC 42001 • SEC REGULATION S-P AMENDMENTS, MARKETING RULE AND EXAM PRIORITIES • FINRA RULES 3110 AND 2210 AND REGULATORY NOTICE 24-09 • NYDFS 23 NYCRR PART 500 AND AI CYBERSECURITY GUIDANCE. THIS SHEET IS GENERAL INFORMATION, NOT LEGAL ADVICE.

TALK TO US

WEB securitybasecamp.com
CALL (949) 330-0899
EMAIL info@securitybasecamp.com

THE SECURE USE OF AI ASSESSMENT GIVES YOU

- ◆ All nine controls scored across four dimensions, with residual risk
- ◆ A complete AI asset inventory, including shadow and vendor-embedded AI
- ◆ Where client data is exposed, and what is technically stopping it
- ◆ A prioritized remediation roadmap with owners and dates
- ◆ The exam folder, assembled — and maintained in EntrustIQ



POWERED BY

EntrustIQ

EntrustIQ is Security Basecamp's sister company, founded in 2025 to productize a decade of assessment methodology. A PDF is accurate the day it is signed and decays from there. EntrustIQ carries the controls, the scoring, and the exam folder as a live structure, so the score moves when your environment does.