



ACF IN ACTION

acfstandard.com

ACF en Action applique l'Agentic Commerce Framework® à des entreprises réelles, à partir de sources publiques. Chaque cas montre qui reste responsable quand l'IA décide.

ACF-CS-001 · Case Study #001 · Édition 2026

Morgan Stanley Wealth Management

How do you govern an assistive AI in a highly regulated financial environment?

CARTE D'IDENTITÉ

SECTOR Finance	COUNTRY USA	TECHNOLOGY GPT-4 / OpenAI	TYPE Assisted AI	AUTONOMY suggestive to co-decision
CONFIDENCE Medium				

OUTILS ACF ILLUSTRÉS DANS CE CAS

● ACF-00	● ACF-01	● ACF-03	● ACF-05	● ACF-02	● ACF-12
----------	----------	----------	----------	----------	----------

● mobilisé · ● à prévoir · 6 outils du Toolkit (4 mobilisés, 2 à prévoir, sur 17).

1. Executive summary

Morgan Stanley is one of the first major Wall Street players to have deployed generative AI at scale with its financial advisors, in partnership with OpenAI. Two tools structure this deployment: **AI @ Morgan Stanley Assistant** (an internal document research assistant) and **AI @ Morgan Stanley Debrief** (a meeting summary tool).

Read through the lens of ACF, this deployment is that of a "**good student**": Morgan Stanley deliberately kept its tools at the **lowest autonomy level** (the AI proposes, the advisor decides), which naturally satisfies the first principle of accountability. The value of ACF here is not corrective, it is **prospective**: it provides the map for the inevitable moment when these tools become **agentic**, that is, when they act without human review at each step. It is at this tipping point that the per-decision accountability layer becomes necessary.

2. Context

Morgan Stanley Wealth Management runs one of the largest networks of financial advisors in the world. The initial constraint: to make a considerable mass of internal research and documentation usable, without ever compromising the compliance of a strictly regulated sector (investment advice).

The deployment is part of a strategic partnership with OpenAI, with Morgan Stanley presented as OpenAI's reference client in wealth management.

3. Public sources used

- Morgan Stanley, official press releases: launch of AI @ Morgan Stanley Debrief; OpenAI partnership milestone; AskResearchGPT.
- OpenAI, "Morgan Stanley uses AI evals to shape the future of financial services".
- CNBC (Sept. 18, 2023): launch of the assistant for financial advisors.
- CNBC (June 26, 2024): rollout of Debrief to wealth management advisors.

4. Reconstructed AI architecture (from public sources)

AI @ Morgan Stanley Assistant - in service since September 2023.

- Nature: internal conversational assistant, backed by GPT-4.
- Knowledge base: ~100,000 internal reports and research documents, curated over months, in a private environment controlled by Morgan Stanley.
- Function: search and synthesis of information for advisors (markets, recommendations, internal procedures). It **answers**, it does not act.
- Adoption: more than 98% of advisor teams use it.

AI @ Morgan Stanley Debrief - pilot March 2024, general rollout mid-2024.

- Nature: summary tool, backed by GPT-4.
- Function: from a meeting recording (Zoom), **with the client's consent**, it generates notes integrated into the CRM and follow-up drafts. The advisor **reviews, adjusts, and validates** before any send.
- Reported gain: ~30 minutes of admin saved per meeting. Adoption: ~98%.

Publicly documented safeguards:

- an **evaluation framework (evals)** testing each use case before deployment (responses scored by advisors and prompt engineers, on accuracy and consistency);
- a **data non-retention** agreement with OpenAI (data is not used to train public ChatGPT);
- **systematic human review**: the advisor validates any output before client use;
- **client consent** for Debrief.

5. ACF mapping

 **Tools used:** ACF-01 (autonomy map) · ACF-04 (agent card) · ACF-03 (agentic constitution).

In ACF language, the architecture reads as follows:

ACF element	In the Morgan Stanley deployment
Named human (accountable)	The financial advisor: it is they who answers for everything that goes to the client. Above: Wealth Management leadership.
Agent	The Assistant (research) and Debrief (summary). They execute, they do not decide.
Platform	GPT-4 / OpenAI, under a non-retention agreement, private Morgan Stanley environment.
Delegated decision	None, in the strong sense. The AI proposes (synthesis, draft), the human decides .
Autonomy level	Suggestive (Assistant) and Co-decision (Debrief: the agent prepares, the advisor validates and sends).
Non-delegable zones	Investment advice and any binding client communication - kept in human hands.
Rule / safeguard	Curated base only (no open web), evals before deployment, human review, client consent, non-retention.
Control third party	Internal compliance function; regulators (SEC, FINRA).
Traceability	Debrief writes to the CRM (trace of notes); the evals are documented. A decisions register in the ACF sense is not required at this autonomy level.

A schema in ACF graphic language accompanies this case (separate file).

6. Analysis

 **Tools used:** ACF-03 (non-delegable zones) · ACF-05 (supervision). **To prepare if autonomy increases:** ACF-08 (register) · ACF-12 (mandate).

Strengths (what Morgan Stanley already does, and does well):

- **Non-delegable zones held.** By keeping advice and client sends under human decision, Morgan Stanley satisfies, without naming it, the principle of non-delegation of accountability: the advisor remains the named human who answers. This is exactly the posture ACF recommends.
- **Pre-deployment discipline.** The evals framework aligns with the ACF logic of prior testing and explicit safeguards.
- **Documentary grounding.** Restricting the Assistant to a curated base limits hallucinations on a terrain where error carries a compliance cost.

Points of attention (where ACF sheds light on what comes next):

- **The agentic threshold has not yet been crossed, but it will be.** The day a tool *acts* (plans, executes, transacts) without review at each step, autonomy rises, and accountability must then be **designed**, not assumed. This is where the need for a designated officer of the delegated decision, a per-agent mandate, and a decisions register appears.
- **Evals measure model quality, not decision accountability.** These are two distinct layers. A well-evaluated model can still produce a decision that no one formally owns, if autonomy increases.
- **Debrief touches client data and the CRM.** Consent and non-retention cover the data; ACF would additionally formalize the trace of the decision (who validated what, on which rule).

7. ACF toolkit used

This grid indicates the tools that the ACF analysis uses (✓) and those to anticipate on the autonomy trajectory (→
SOON). It does not mean the company uses ACF: it shows which instruments the framework would bring.

Tool	Status	Role in this case
ACF-00 · Sovereignty Score	✓	Assess maturity (low to moderate exposure)
ACF-01 · Autonomy map	✓	Position each use (suggestive / co-decision)
ACF-03 · Agentic constitution	✓	Formalize the non-delegable zones already held
ACF-05 · Supervision	✓	Link evals to continuous monitoring
ACF-06 · Kill switch	→ SOON	To anticipate if autonomy increases
ACF-08 · Decisions register	→ SOON	As soon as an agent acts without review
ACF-12 · Agent mandate	→ SOON	Bound each future agent

8. Governance questions (what an AI committee should ask itself)

1. For each tool, at which autonomy level is it today, and who decided that level?
2. The day an agent acts without review at each step, **who will answer for it by name?**
3. Which decisions do we declare non-delegable, explicitly and signed?
4. Do we have a trace, reconstructable and enforceable, of who validated what and on which rule?
5. Do our evaluations test model quality, decision accountability, or both?

9. General lessons

Even a prudent and praised deployment gains from making **explicit** what it does implicitly well. Morgan Stanley holds its non-delegable zones by design choice; ACF makes them nameable, signable, and auditable, and above all it provides the map for the move to higher autonomy, which is the trajectory of the entire sector.

The transferable lesson: good AI governance is not about constraining tools, but about knowing, at each notch of autonomy, who remains accountable.

10. What ACF would do (reference design)

Prescriptive and non-critical formulation: here is how an organization would design such a system directly with the framework, whatever its sector.

If an organization wished to design this type of system directly with ACF, it could in particular:

- **map** each use on the autonomy scale (ACF-01) and explicitly set the authorized level, decision by decision;
- **declare** its non-delegable zones in a signed agentic constitution (ACF-03) - here: investment advice and client sends;
- **designate** an officer of the delegated decision (DDAO), activatable as soon as an agent switches to autonomous mode;
- **document** a decisions register (ACF-08) linking each agent action to the rule that authorized it and to the human who answers for it;
- **define** a stop mechanism (ACF-06) and its triggers;
- **link** model evaluations to continuous accountability supervision (ACF-05) - two distinct layers.

None of this is a reproach addressed to Morgan Stanley: it is the way ACF would make **explicit, signable, and auditable** what a prudent deployment already does implicitly - and would equip it for the rise in autonomy.

11. Similar cases (transferability)

This case is specific neither to Morgan Stanley nor to banking. It is particularly transferable to organizations that combine **many documents, experts, and a strong compliance constraint**:

-  banks and wealth management
-  insurance
-  consulting firms
-  law firms
-  healthcare (hospitals, laboratories)
-  audit and accounting

In all these sectors, the same pattern repeats: an AI that synthesizes and drafts from internal knowledge, a professional who remains accountable, and a boundary - advice, the act, the binding decision - never to be crossed without a named human. This is what makes this document a **reference case**: it illustrates a concept, not just a company.

12. Confidence of the assessment

High. The case rests on abundant, first-hand public documentation: Morgan Stanley official press releases, an OpenAI publication, and reference business press. The reported facts are therefore solidly established and verifiable by the reader. The ACF analysis draws on these public facts; only certain internal implementation details are not public, without bearing on the governance reading proposed here.

Teacher guide

Pedagogical objective. Convey the difference between *assisted* AI and *agentic* AI, and why accountability must be designed before increasing autonomy.

Session outline (90 min).

1. (15 min) Reading of the case, without the Analysis part.
2. (20 min) In groups: map the deployment in ACF language (who executes, who answers, which non-delegable zones).
3. (20 min) Debate: "Is Morgan Stanley compliant with the non-delegation principle?" (expected answer: yes, through keeping autonomy low).
4. (20 min) Projection: "We switch Debrief to autonomous mode - it sends follow-ups without review. What changes?" Bring out the need for a designated officer and a register.
5. (15 min) Synthesis by the teacher.

Possible exam questions. See the "Governance questions" section.

Reminder: grading students' work is a pedagogical act of the teacher. This guide provides the reference analysis, not automatic grading.

Associated lab (ACF-LAB-001)

From this case, the student produces:

- the Human-AI-Agent mapping of the two tools (card ACF-O1);
- the criticality matrix of three hypothetical decisions if Debrief became autonomous (card ACF-O2);
- an agentic constitution for an "autonomous Debrief" (card ACF-O3): non-delegable zones, caps, escalation conditions;
- the definition of a kill switch (card ACF-O6) and its trigger.

Appendix - Completed ACF cards

The official Toolkit cards used in this case, filled in with its public data. For illustration: an organization of the same profile can reuse and adapt them.



OBJECTIVE

Assess the governance maturity of the deployment, from public sources only. Qualitative assessment: a numeric score would require internal data.

1 THE SIX DIMENSIONS

Dimension	Assessment (public sources)
Decision autonomy	Sound - the human decides, the AI supports
Transparency	Good - curated base, documented evals
Non-delegable zones	Held - advice stays human
Supervision	Continuous - evals + systematic review
Decision traceability	Partial - CRM notes, no dedicated register
Living governance	To formalize - DDAO not designated

2 VERDICT



Low to moderate exposure. Low autonomy and the non-delegable zones held limit the risk of dependence.

POINTS OF ATTENTION

- Formalize decision traceability (register)
- Designate a DDAO before any rise in autonomy

On an indicative scale, this profile corresponds to a **moderate exposure**: the foundations are healthy, the room for progress is on formal governance, not on the technique.



OBJECTIVE Map the decisions entrusted to Morgan Stanley's AI tools and set, for each, the autonomy level. Filled from public sources.

1 AGENTIC MATURITY SCALE

0

Classic automation

Fixed rules, no autonomy.

1

Assisted agents

The agent proposes, the human decides.

2

Governed agents

The agent acts within a frame, under supervision.

TARGET ~80%

3

Supervised autonomy

The agent decides alone; control after the fact.

2 YOUR KEY DECISIONS (filled)

A

Document research

STAKE

Exploit ~100,000 curated documents

IMPACT

Time saved, sourced answers

DEADLINE

Continuous (Assistant)

AUTONOMY

Suggestive

B

Meeting summary

STAKE

CRM notes and follow-up drafts (Debrief)

IMPACT

~30 min / meeting

DEADLINE

Post-meeting

AUTONOMY

Co-decision

C

Investment advice

STAKE

Binding recommendation to the client

IMPACT

Commits the firm's accountability

DEADLINE

-

AUTONOMY

NON-DELEGABLE

Reading: three decisions, three levels. Two at low autonomy, the third non-delegable. That is what keeps accountability naturally held.



OBJECTIVE

Define the agent's constitution: identity, standing principles and forbidden zones. Card filled with public case data.

1 AGENT IDENTITY & MISSION

AGENT NAME

Assistant & Debrief (AI @ Morgan Stanley)

ACCOUNTABLE (DDAO)

To be designated (de facto: advisor + WM leadership)

PRIMARY MISSION

Document research and meeting summary, in support of the advisor

SCOPE

Curated internal research base; meeting notes (with consent)

2 CORE PRINCIPLES

- 1 Respect SEC / FINRA compliance in all circumstances
- 2 Rely only on the curated base (no open web)
- 3 Submit every output to human review before client use
- 4 Escalate rather than decide on advice

EXAMPLE PRINCIPLES

- Respect the legal and sector framework
- Protect data and confidentiality
- Stay transparent, every decision traceable
- When in doubt, escalate rather than decide

3 FORBIDDEN ZONES (ABSOLUTE)

- 1 Issue investment advice autonomously
- 2 Send a client communication without human validation
- 3 Process data outside the non-retention environment

EXAMPLE PROHIBITIONS

- X Commit spending above a threshold without validation
- X Share personal data with a third party
- X Change its own rules or scope on its own

Constitution filled from publicly documented practices. The DDAO remains to be formally designated: it is the main governance task before any rise in autonomy.



OBJECTIVE

Link model evaluations to continuous accountability monitoring. Card filled with the publicly known indicators of the case.

1 SUPERVISION INDICATORS

Indicator	Value / target (public sources)
Adoption	98% of advisor teams
Answer quality	evaluated continuously (evals framework)
Human review	100% of client-facing outputs
Client consent	required for Debrief
Escalation / error rate	to be tracked (not public)

2 MECHANISM

- 1 Evals framework: test each use case before deployment
- 2 AI Act compliance matrix: human oversight (art. 14), logging (art. 12)
- 3 Supervision officer: to be designated (DDAO)
- 4 Periodic review of the delegated scope

Evals measure the quality of the **model**. To cover accountability for the **decision**, link this monitoring to the register (ACF-08) as soon as autonomy rises.

Disclaimer

This study is carried out **exclusively from public information**. It constitutes **neither an audit, nor a certification, nor an official evaluation** of Morgan Stanley. It applies the ACF framework for pedagogical and illustrative purposes. No internal, confidential, or non-public data was used. The trademarks cited belong to their respective owners.

Sources

- CNBC, *Morgan Stanley kicks off generative AI era on Wall Street with assistant for financial advisors* (Sept. 18, 2023) - <https://www.cnbc.com/2023/09/18/morgan-stanley-chatgpt-financial-advisors.html>
- CNBC, *Morgan Stanley wealth advisors are about to get an OpenAI-powered assistant to do their grunt work* (June 26, 2024) - <https://www.cnbc.com/2024/06/26/morgan-stanley-openai-powered-assistant-for-wealth-advisors.html>
- OpenAI, *Morgan Stanley uses AI evals to shape the future of financial services* - <https://openai.com/index/morgan-stanley/>
- Morgan Stanley, *Launch of AI @ Morgan Stanley Debrief* - <https://www.morganstanley.com/press-releases/ai-at-morgan-stanley-debrief-launch>
- Morgan Stanley, *Key Milestone in Innovation Journey with OpenAI* - <https://www.morganstanley.com/press-releases/key-milestone-in-innovation-journey-with-openai>
- Morgan Stanley, *Morgan Stanley Research Announces AskResearchGPT* - <https://www.morganstanley.com/press-releases/morgan-stanley-research-announces-askresearchgpt>

ACF Case Study ACF-CS-001 · v1 · working document, to be reviewed before any public distribution.