



ACF EN 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-025 · Case Study #025 · Édition 2026

Morgan Stanley (la supervision continue)

Comment faire vivre un agent IA dans la durée sans qu'il dérive ?

CARTE D'IDENTITÉ

SECTEUR	PAYS	TECHNOLOGIE	TYPE	AUTONOMIE	CONFIANCE
Finance	USA	Supervision IA	IA assistée	supervisé	Moyenne

OUTILS ACF ILLUSTRÉS DANS CE CAS

● ACF-01	● ACF-05	● ACF-06	● ACF-08	● ACF-00	● ACF-02
----------	----------	----------	----------	----------	----------

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

1. Résumé exécutif

Morgan Stanley déploie à partir de 2023 un assistant IA (adossé à GPT-4) pour ses quelque 16 000 conseillers en gestion de patrimoine. L'agent aide à retrouver l'information dans une vaste base documentaire interne, puis à préparer des comptes rendus de réunion. Plutôt que de le lancer et de l'oublier, Morgan Stanley l'entoure d'un dispositif de supervision continue : chaque cas d'usage est testé (evals) avant déploiement, la production passe par une relecture humaine avant tout usage engageant, et le tout est journalisé dans un cadre de conformité. Résultat : l'agent est largement adopté (de l'ordre de 98 % des équipes concernées), tout en restant sous contrôle.

Lu au prisme de l'ACF, ce cas est le pendant « **référence** » de DPD : les deux traitent le même objet, la **dérive et la supervision d'un agent dans la durée**, par des chemins opposés. Là où DPD a laissé un changement casser ses garde-fous sans détection, Morgan Stanley conçoit la confiance comme un **dispositif** qui teste, relit et journalise en continu. La valeur de l'ACF est ici **prospective** : elle nomme ce que ce déploiement fait bien et fournit le langage pour l'étendre, quel que soit le secteur.

Note : Morgan Stanley sert aussi de référence au kit « Responsabilité » (autonomie basse, zones non déléguables tenues). Ici, l'angle est différent : la supervision continue qui évite la dérive. Le même déploiement éclaire deux objets ACF distincts.

2. Contexte

Morgan Stanley Wealth Management gère l'un des plus grands réseaux de conseillers financiers au monde, dans un secteur strictement régulé. Comme tout agent déployé à grande échelle, cet assistant évolue dans la durée : nouveaux cas d'usage, nouvelles capacités, mises à jour. La question que le cas pose est celle de la longévité : comment un agent reste-t-il sûr non pas au lancement, mais à chaque évolution ?

3. Sources publiques utilisées

- Morgan Stanley, communiqués sur le déploiement de l'assistant IA (2023-2024).
- OpenAI, page consacrée au déploiement de Morgan Stanley.
- Reprises presse (CNBC) rapportées.

4. Architecture IA reconstituée (à partir de sources publiques)

Assistant IA Morgan Stanley (en service depuis 2023, étendu en 2024).

- **Nature** : assistant conversationnel interne, adossé à GPT-4, réservé aux conseillers.
- **Fonction** : retrouver et synthétiser l'information dans une base documentaire interne, puis préparer des comptes rendus de réunion. Il **produit des éléments** que le conseiller utilise ; il n'engage pas le client seul.
- **Adoption** : de l'ordre de 98 % des équipes concernées.

Le dispositif de supervision (ce que les faits révèlent) :

- **Evals avant déploiement** : chaque cas d'usage est testé avant d'être mis en service.
- **Relecture humaine** : la production de l'agent est relue par un humain avant tout usage engageant vis-à-vis du client.
- **Conformité** : journalisation et supervision humaine, alignées sur les exigences réglementaires.
- **Périmètre tenu** : l'agent assiste ; le conseil et l'envoi client restent des actes humains.

La logique de gouvernance. Le point clé n'est pas que l'agent soit « parfait », mais qu'il soit **supervisé dans la durée** : évalué avant chaque déploiement, relu, journalisé. Un changement ou une nouvelle capacité passe par les evals avant la mise en service, ce qui rend une dérive détectable **avant** qu'elle n'atteigne le client. La confiance repose sur un dispositif, pas sur l'espoir que l'agent ne déviera jamais.

5. Cartographie ACF

 **Outils mobilisés** : ACF-05 (supervision continue) · ACF-03 (constitution agentique) · ACF-01 (cartographie).

En langage ACF, le déploiement se lit ainsi :

Élément ACF	Dans le déploiement Morgan Stanley
Humain nommé (responsable)	Le conseiller financier pour chaque usage client ; au-dessus, la direction Wealth Management.
Agent	L'assistant IA (recherche documentaire, préparation de comptes rendus). Il produit, il n'engage pas seul.
Plateforme	GPT-4 / OpenAI, dans un environnement contrôlé par Morgan Stanley.
Décision déléguée	Aucune décision engageante au sens fort. L'agent prépare , l'humain relit et tranche .
Niveau d'autonomie	Suggestif à co-décision : l'agent propose, le conseiller valide avant tout usage client.
Zones non délégables	Le conseil et l'envoi client engageant, maintenus entre des mains humaines.
Règle / garde-fou	Evals avant déploiement, relecture humaine avant usage engageant, journalisation, périmètre borné.
Tiers de contrôle	Fonction conformité interne ; régulateurs du secteur.
Traçabilité	Journalisation dans un cadre de conformité ; les evals sont documentées et rejouées à chaque évolution.

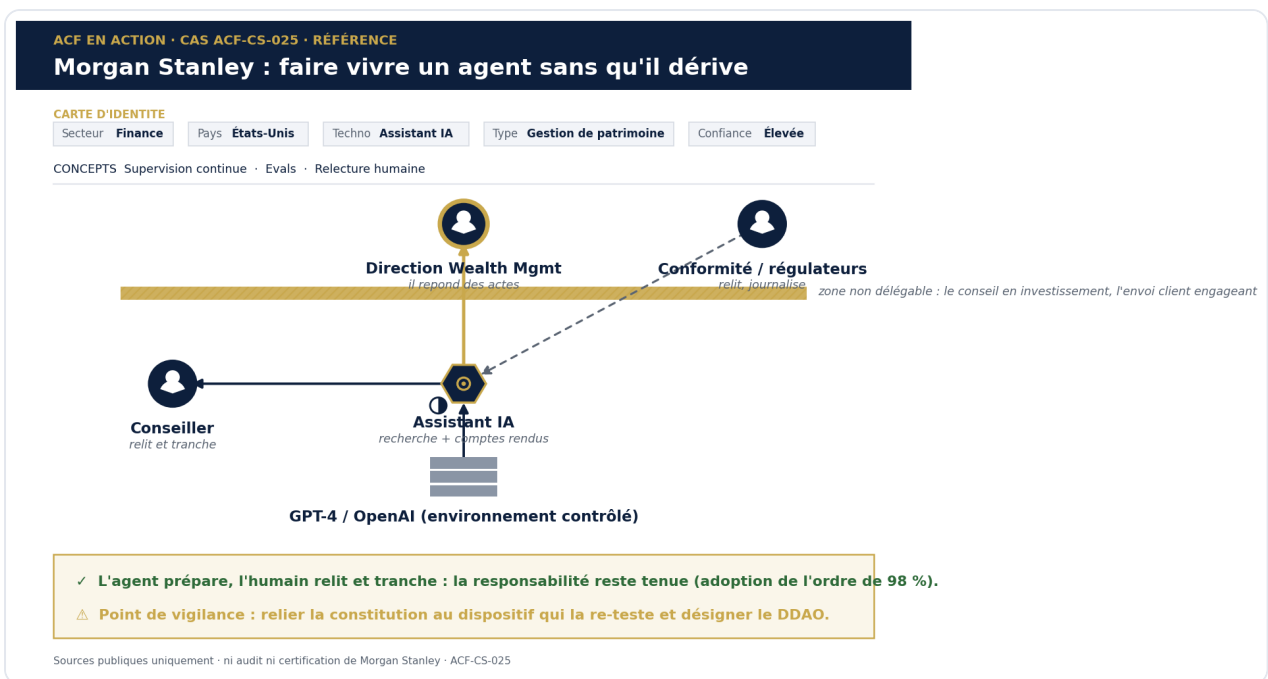
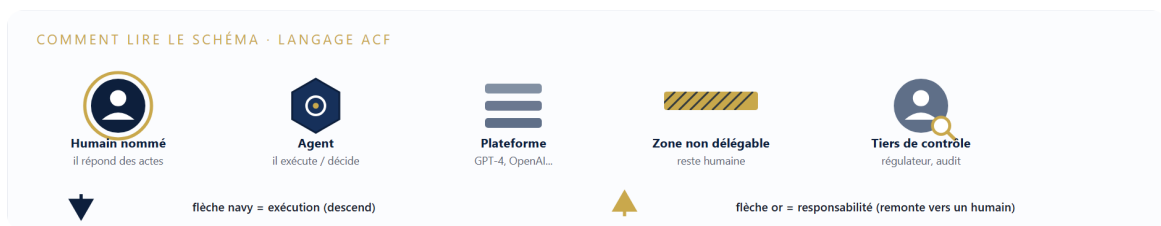


Figure 1 - Morgan Stanley (la supervision continue) en langage ACF (Language Card).



6. Analyse

Outils mobilisés : ACF-05 (evals + relecture + journal) · ACF-03 (constitution vérifiée) · ACF-01 (cartographie).

Ce que ce déploiement fait bien (et que l'ACF nomme) :

- **La supervision est un dispositif, pas un événement.** Morgan Stanley ne parie pas sur un agent infaillible : il le teste, le relit et le journalise en continu. La différence avec DPD n'est pas la qualité du modèle, c'est le dispositif de supervision autour de lui.
- **La dérive est détectable avant le client.** Un écart apparaît d'abord dans les evals, avant la mise en service. Tester chaque cas d'usage protège le client d'un comportement non vérifié.
- **La production reste filtrée.** La relecture humaine avant tout usage engageant place un humain nommé sur le chemin, exactement là où la responsabilité doit rester.

Ce que l'ACF éclaire pour la suite :

- **La constitution doit être reliée au dispositif qui la teste.** Écrire des règles ne suffit pas ; elles ne valent que si elles sont **vérifiées** après chaque évolution, pas seulement au lancement.
- **La bonne gouvernance ne vise pas la perfection.** Elle vise un dispositif qui **détecte et corrige** la dérive à temps. C'est un principe transférable, pas un privilège de Morgan Stanley.
- **La montée en autonomie changera la donne.** Si un agent agit un jour sans relecture à chaque acte, la supervision devra être conçue comme telle : evals continues, seuils, journal des décisions.

7. Toolkit ACF mobilisé

Cette grille indique les outils que l'analyse ACF mobilise (✓) et ceux à prévoir sur la trajectoire (→
SOON). Elle ne signifie pas que l'entreprise utilise l'ACF : elle montre quels instruments le référentiel apporterait.

Outil	Statut	Rôle dans ce cas
ACF-05 · Supervision continue	✓	Relier evals, relecture humaine et journalisation en un dispositif permanent
ACF-03 · Constitution agentique	✓	Formaliser le périmètre et garantir son re-test à chaque changement
ACF-01 · Cartographie	✓	Situer l'agent, le responsable et le niveau d'autonomie
ACF-08 · Registre des décisions	→ SOON	Dès qu'un agent agit sans relecture à chaque acte
ACF-06 · Kill switch	→ SOON	À prévoir si l'autonomie augmente

8. Questions de gouvernance (ce qu'un comité IA devrait se poser)

1. Pourquoi tester chaque cas d'usage avant déploiement protège-t-il de la dérive ?
2. La supervision est-elle un événement (au lancement) ou un travail continu ?
3. Un agent doit-il être « parfait », ou faut-il un dispositif qui détecte et corrige à temps ?
4. Nos règles de constitution sont-elles reliées à un dispositif qui les re-teste après chaque évolution ?
5. Qui relit la production de l'agent avant tout usage engageant vis-à-vis du client ?

9. Enseignements généraux

Un agent n'est pas sûr parce qu'il « marche bien » : il est sûr parce qu'il est supervisé dans la durée. Morgan Stanley ne mise pas sur un modèle infallible, il installe un dispositif qui teste avant chaque déploiement, relit avant chaque usage engageant et journalise ce qui est produit. C'est ce dispositif, pas la qualité brute du modèle, qui rend la dérive détectable avant qu'elle n'atteigne le client.

La leçon transférable : la confiance dans un agent repose sur un dispositif de supervision continue, pas sur l'espoir qu'il ne dévierait jamais. Superviser est un travail permanent, pas une case cochée au lancement.

10. Ce que ferait l'ACF (conception de référence)

Formulation prescriptive et non critique : voici comment une organisation concevrait un tel système directement avec le référentiel, quel que soit son secteur.

Si une organisation souhaitait concevoir ce type d'agent directement avec l'ACF, elle pourrait notamment :

- **cartographier** chaque usage sur l'échelle d'autonomie (ACF-01) et fixer le niveau autorisé, en gardant le conseil et l'envoi client hors du périmètre de l'agent ;
- **déclarer** dans une constitution agentique signée (ACF-03) le périmètre et le passage obligé par la relecture humaine avant tout usage engageant ;
- **installer** un dispositif de supervision continue (ACF-05) : evals avant chaque déploiement, relecture humaine, journalisation, re-test des règles à chaque changement ;
- **relier** explicitement la constitution au dispositif qui la vérifie, pour qu'aucune évolution ne l'efface silencieusement ;
- **préparer** un registre des décisions (ACF-08) et un dispositif d'arrêt (ACF-06) pour le jour où l'autonomie augmentera.

Rien de tout cela n'est un reproche adressé à Morgan Stanley : c'est la manière dont l'ACF rendrait **explicite, signable et auditable** la supervision continue qu'un déploiement mûr pratique déjà, et l'outillerait pour la montée en autonomie.

11. Cas similaires (transférabilité)

Ce cas n'est propre ni à Morgan Stanley ni à la finance. Il est particulièrement transférable aux organisations qui font **vivre un agent dans la durée sous contrainte de fiabilité** :

-  banques et gestion de patrimoine
-  assurances
-  santé (hôpitaux, laboratoires)
-  industrie et énergie (agents d'assistance métier)
-  services publics et administrations
-  éditeurs de logiciels exploitant des agents évolutifs

Partout, le même schéma : un agent qui produit à grande échelle, une évolution permanente, et une supervision qui teste, relit et journalise à chaque changement. C'est ce qui fait de ce document un **cas de référence** : il illustre un concept, pas seulement une entreprise.

12. Confiance de l'évaluation

Élevée. Le cas repose sur une documentation publique de première main : communiqués de Morgan Stanley, publication d'OpenAI et reprises de presse économique de référence. Les faits rapportés sont donc solidement établis et vérifiables par le lecteur. L'analyse ACF s'appuie sur ces faits publics ; seuls certains détails internes du dispositif de supervision ne sont pas publics, sans incidence sur la lecture de gouvernance proposée ici.

Guide enseignant

Objectif pédagogique. Faire comprendre que la supervision d'un agent est un travail continu, et que la confiance repose sur un dispositif (evals, relecture, journal), pas sur la qualité brute du modèle.

Déroulé de séance (90 min).

1. (15 min) Lecture des **deux** cas du kit : DPD (incident) et Morgan Stanley (référence), sans les analyses.
2. (20 min) En groupes : décrire le dispositif de supervision continue en langage ACF (fiche ACF-05).
3. (20 min) Atelier : écrire la constitution et le re-test à chaque changement (fiche ACF-03).
4. (20 min) Comparaison DPD / Morgan Stanley : changement non contrôlé contre evals avant déploiement.
5. (15 min) Synthèse : superviser un agent est un travail continu, qui accompagne chaque évolution.

Questions d'examen possibles. Voir la section « Questions de gouvernance ».

Rappel : la correction du travail des étudiants est un acte pédagogique du professeur. Le présent guide fournit l'analyse de référence, pas une notation automatique.

Lab associé (ACF-LAB-025)

À partir de ce cas, l'étudiant produit :

- le dispositif de supervision continue (fiche ACF-05) : evals avant déploiement, relecture humaine, journalisation ;
- une Constitution agentique (fiche ACF-03) : périmètre, passage par la relecture, re-test à chaque changement ;
- la cartographie de l'agent (fiche ACF-01) : agent, responsable, niveau d'autonomie ;
- une comparaison écrite DPD / Morgan Stanley centrée sur la supervision.

Annexe - Fiches ACF complétées

Les fiches officielles du Toolkit mobilisées dans ce cas, remplies avec ses données publiques. À titre illustratif : une organisation du même profil peut les reprendre et les adapter.



OBJECTIF

Cartographier les usages de l'assistant IA de Morgan Stanley et fixer, pour chacun, le niveau d'autonomie. L'agent prépare, l'humain relit et tranche. Rempli à partir des sources publiques.

1 ÉCHELLE DE MATURITÉ AGENTIQUE

0

Automatisation classique

Règles fixes, aucune autonomie.

1

Agents assistés

L'agent propose, l'humain décide.

2

Agents gouvernés

L'agent agit dans un cadre, sous supervision.

CIBLE ~80 %

3

Autonomie supervisée

L'agent décide seul ; contrôle a posteriori.

2 VOS DÉCISIONS CLÉS (rempli)

A

Recherche documentaire

ENJEU

Retrouver l'information dans la base curée

IMPACT

Réversible, réponses sourcées

ÉCHÉANCE

Continu

AUTONOMIE

Suggestif

B

Compte rendu de réunion

ENJEU

Préparer notes et brouillons de suivi

IMPACT

Produit relu avant tout usage

ÉCHÉANCE

Post-réunion

AUTONOMIE

Co-décision (relecture humaine)

C

Conseil en investissement, envoi client engageant

ENJEU

Recommander, engager la firme

IMPACT

Engage la responsabilité

ÉCHÉANCE

-

AUTONOMIE

NON DÉLÉGABLE

Lecture : deux usages en autonomie basse, le conseil non délégable. C'est cette carte, **re-testée à chaque évolution**, qui maintient la dérive détectable avant le client.



OBJECTIF

Relier les évaluations de modèle, la relecture humaine et la journalisation en un dispositif permanent. Fiche remplie avec les éléments publiquement connus du cas.

1 INDICATEURS DE SUPERVISION

Indicateur	Valeur / dispositif (sources publiques)
Evals avant déploiement	Chaque cas d'usage testé avant mise en service
Relecture humaine	Avant tout usage engageant vis-à-vis du client
Journalisation	Dans un cadre de conformité (supervision humaine)
Adoption	De l'ordre de 98 % des équipes concernées
Re-test à chaque évolution	Evals rejouées avant toute nouvelle capacité

2 DISPOSITIF

- 1 Evals : test de chaque cas d'usage avant déploiement et après chaque changement
- 2 Relecture humaine systématique avant tout usage engageant
- 3 Journalisation alignée sur les exigences de conformité
- 4 Responsable de supervision : le DDAO, à désigner formellement

La supervision est un **dispositif, pas un événement**. Relier les evals au registre (ACF-08) dès que l'autonomie monte, pour couvrir la responsabilité de la décision, pas seulement la qualité du modèle.



OBJECTIF

Définir les signaux qui bloquent un cas d'usage ou basculent vers un humain. Fiche remplie en dispositif prospectif, à armer avec la montée en autonomie.

1 DÉCLENCHEURS

Signal	Seuil	Action
Échec d'une eval après changement	Détecté	Blocage de la mise en service du cas d'usage
Écart récurrent en relecture humaine	Au-dessus du seuil	Suspension du cas d'usage, revue
Usage hors périmètre (conseil, envoi engageant)	Détecté	Blocage et escalade humaine

2 PROCÉDURE

- 1 Qui décide l'arrêt : le DDAO
- 2 Comment : blocage du cas d'usage concerné, relecture renforcée, escalade humaine
- 3 Test régulier du dispositif, avant chaque montée en autonomie
- 4 Journalisation de chaque blocage (ACF-08)

Ici le dispositif est **prospectif** : tant que l'humain relit chaque sortie engageante, l'arrêt reste manuel. Il devient nécessaire dès que l'agent agit sans relecture à chaque acte.



OBJECTIF

Consigner les productions engageantes et leur relecture pour en répondre. Fiche remplie avec les champs pertinents pour le cas.

1 CHAMPS À TRACER

Champ à tracer	Exemple, cas Morgan Stanley
Horodatage	2023-2024 (déploiement puis extension)
Usage / session	Recherche documentaire ou compte rendu de réunion
Production de l'agent	Réponse sourcée / brouillon de suivi
Relecture humaine	Effectuée avant tout usage client engageant
Statut	Validé par le conseiller, relu en conformité

2 USAGE DU REGISTRE

- 1 Documenter que chaque sortie engageante a été relue
- 2 Relier les evals à la trace des décisions à mesure que l'autonomie monte
- 3 Établir qui répond de la production : le conseiller, puis la direction
- 4 Fournir une preuve opposable en cas de contrôle réglementaire

Les evals mesurent le **modèle** ; le registre couvre la **décision**. Les relier dès la montée en autonomie ferme le dernier angle mort de la supervision.

Disclaimer

Cette étude est réalisée **exclusivement à partir d'informations publiques**. Elle ne constitue **ni un audit, ni une certification, ni une évaluation officielle** de Morgan Stanley. Elle applique le référentiel ACF à des fins pédagogiques et illustratives. Aucune donnée interne, confidentielle ou non publique n'a été utilisée. Les marques citées appartiennent à leurs détenteurs respectifs.

Sources

- Morgan Stanley, communiqués sur le déploiement de l'assistant IA (2023-2024)
- OpenAI, page consacrée au déploiement de Morgan Stanley
- CNBC, reprises presse rapportées sur le déploiement de l'assistant IA

ACF Case Study ACF-CS-025 · v1 · document de travail, à relire avant toute diffusion publique.