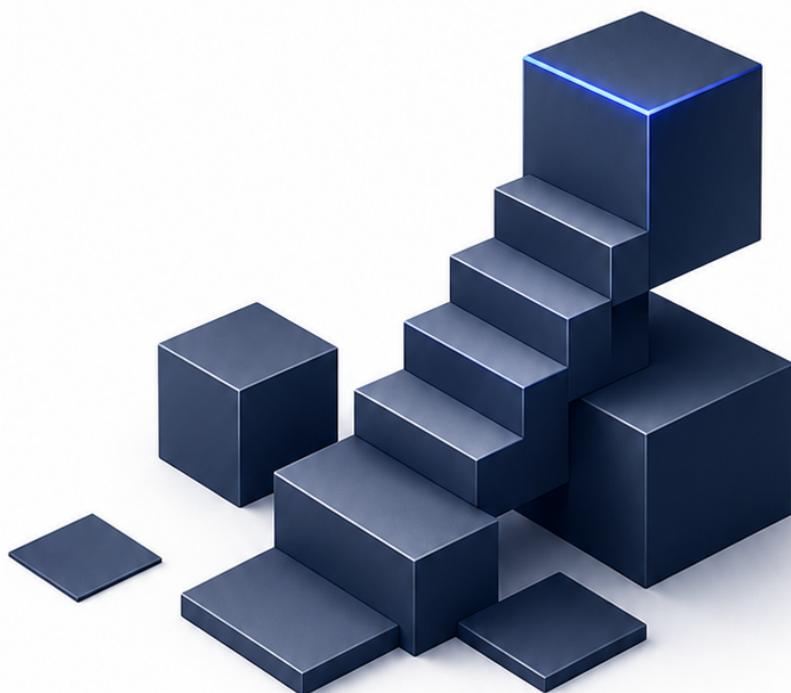


LIVRE BLANC · FRAMEWORK STAIR

LA PREUVE ALGORITHMIQUE

*Adopter l'IA vite ne suffit plus.
La gouverner, si.*

Rénald Féré



Sommaire

Chapitre 1 : L'IA a quitté le pilote

Chapitre 2 : Le mythe de la vitesse

Chapitre 3 : La dette algorithmique

Chapitre 4 : RAG et MCP : les nouveaux angles morts

Chapitre 5 : Le mur réglementaire

Chapitre 6 : Sortir des silos

Chapitre 7 : Intégrer l'IA dans les nouveaux process sans perdre la maîtrise

Chapitre 8 : Le cadre de la preuve algorithmique

Chapitre 9 : Qui gouverne quoi

Chapitre 10 : La feuille de route 90 jours pour un COMEX

Chapitre 11 : De la preuve à l'avantage concurrentiel

À propos de l'auteur

Chapitre 1 : L'IA a quitté le pilote

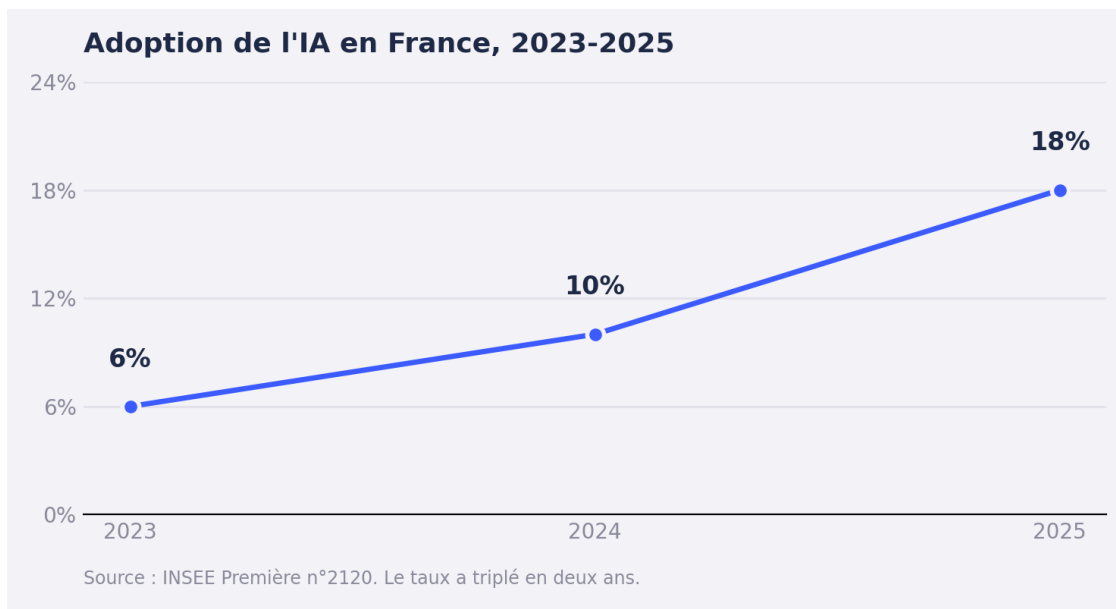
PARTIE 1 : CONSTAT

Pendant deux ans, la phrase que j'ai le plus entendue dans les comités de direction a été : « nous testons l'IA générative ». En 2026, cette phrase est devenue fausse pour la plupart des organisations, et le fait qu'elle continue d'être prononcée est en soi un problème.

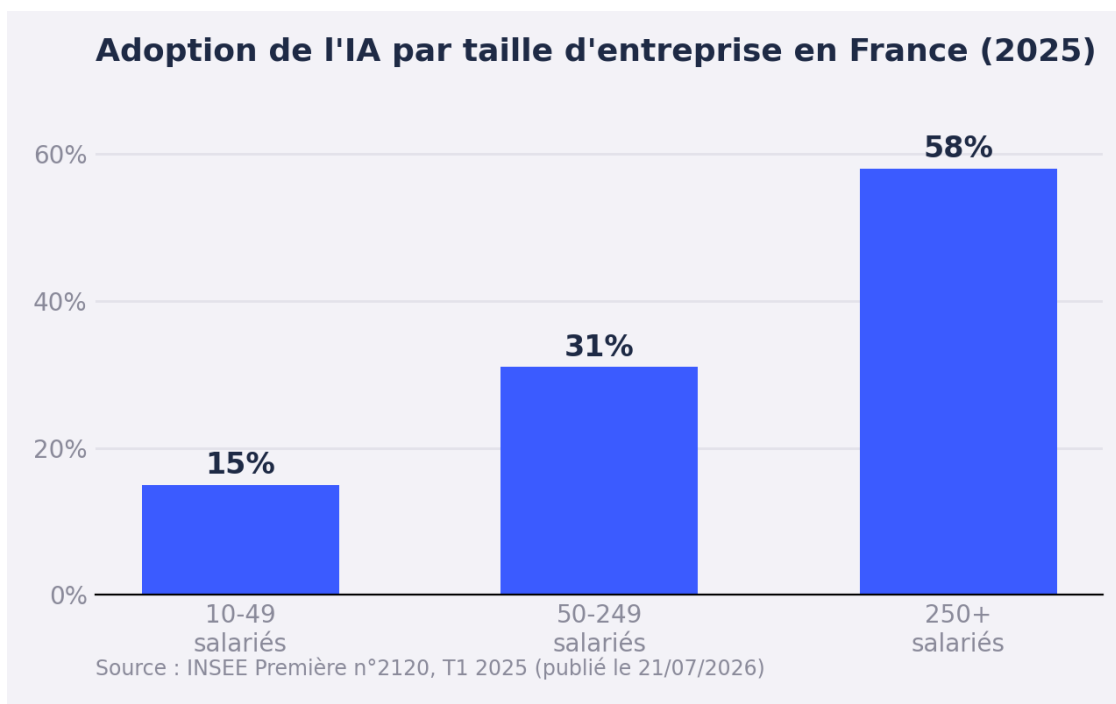
L'IA générative n'est plus un pilote isolé, cantonné à une équipe innovation ou à quelques prompts dans un outil de rédaction. Elle est entrée dans le système d'information. Elle traite des tickets client, résume des contrats, alimente des CRM, génère des campagnes, et de plus en plus, elle agit, via des agents connectés à des systèmes d'entreprise, plutôt que de simplement répondre à une question.

Ce que disent les chiffres

En France, la part des entreprises utilisant au moins une technologie d'intelligence artificielle a été multipliée par trois en deux ans : 6 % en 2023, 18 % en 2025 [1]. Le taux d'adoption reste très corrélé à la taille de l'organisation, 15 % pour les entreprises de 10 à 49 salariés, 31 % pour celles de 50 à 249, 58 % pour celles de 250 salariés et plus [1], mais la trajectoire est la même partout : ascendante, et de plus en plus rapide.



Progression de l'adoption de l'IA en France, 2023-2025



Adoption de l'IA par taille d'entreprise en France (2025)

À l'échelle mondiale, le mouvement est encore plus net : 88 % des organisations utilisent désormais l'IA dans au moins une fonction métier, contre 78 % un an plus tôt [2]. Plus révélateur encore : 62 % des entreprises expérimentent ou déploient déjà des agents IA, des systèmes capables d'agir, pas seulement de répondre, mais seulement 23 % parviennent à les faire fonctionner à l'échelle [2]. C'est cet écart, entre l'expérimentation généralisée et la mise à l'échelle rare, qui définit la situation réelle de 2026.

Côté dirigeants de PME et ETI françaises, le sujet a changé de statut : 58 % jugent désormais l'IA « critique pour la survie » de leur entreprise à horizon 3-5 ans, et 43 % ont formalisé une stratégie IA [3]. Ce n'est plus un sujet d'innovation. C'est un sujet de compétitivité, au même titre qu'un système d'information ou une politique commerciale.

Le vrai changement n'est pas l'usage, c'est l'infiltration

Ce que ces chiffres ne montrent pas, et que je constate dans chacune de mes missions de direction digitale, c'est où l'IA générative s'est installée. Pas seulement dans les outils dédiés que la DSI a validés, mais dans les interstices du système d'information : un connecteur ajouté à un CRM par une équipe commerciale, un assistant branché sur une base documentaire interne par une équipe marketing, un agent qui automatise une tâche RH sans être passé par un comité d'architecture.

C'est un écart que je retrouve dans presque toutes les organisations que j'accompagne : entre l'IA « officiellement déployée » et l'IA « réellement utilisée » dans les process, il y a souvent plusieurs mois et plusieurs outils de décalage. Quand je fais l'inventaire des usages réels au démarrage d'une mission, la cartographie que me remet la DSI est systématiquement plus courte que la réalité constatée dans les équipes.

Cette infiltration n'est ni bonne ni mauvaise en soi. C'est simplement la façon dont les technologies utiles se répandent dans une organisation : par la marge, avant d'atteindre le centre. Le problème n'est pas que l'IA se diffuse vite. Le problème est que la plupart des organisations n'ont pas encore construit les moyens de savoir, à un instant donné, où elle se trouve, ce qu'elle a produit, et sur quelles données elle s'est appuyée.

Pourquoi ce chapitre compte

Le reste de ce livre blanc part d'un principe que je pose d'emblée : on ne peut pas gouverner ce qu'on ne voit pas. Avant de parler de risque (Partie 2), de transformation (Partie 3) ou de gouvernance (Partie 4), il faut d'abord accepter ce constat, inconfortable pour beaucoup de comités de direction : l'IA n'est plus un projet piloté depuis le centre. C'est une infrastructure diffuse, déjà largement déployée, et la vraie question stratégique de 2026 n'est plus « faut-il adopter l'IA générative ? », la réponse est déjà oui, à 88 % dans le monde et en croissance rapide en France, mais « savons-nous ce qu'elle fait, dans notre organisation, en ce moment même ? »

Dans la grande majorité des organisations où je pose cette question, la réponse honnête est non. C'est le point de départ de tout ce qui suit.

Sources

[1] INSEE Première n°2120, « Les technologies de l'information et de la communication dans les entreprises en 2025 », publié le 21 juillet 2026.

[2] McKinsey QuantumBlack, « The state of AI in 2025 », novembre 2025.

[3] Bpifrance Le Lab, « L'IA dans les PME et ETI françaises : une révolution tranquille », juin 2025.

Chapitre 2 : Le mythe de la vitesse

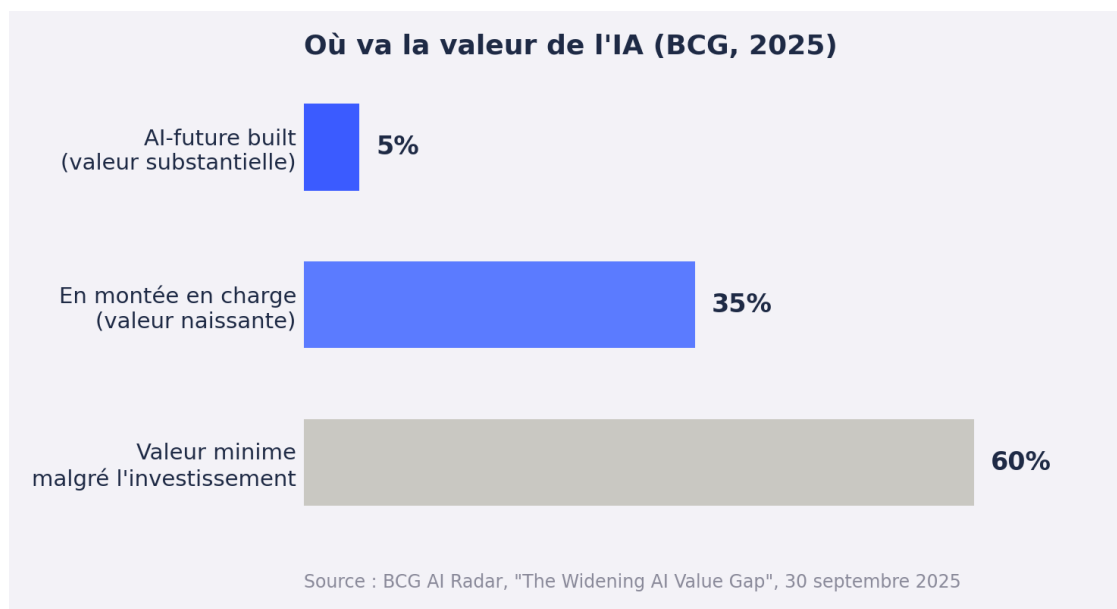
PARTIE 1 : CONSTAT

Dans la plupart des comités de direction où j'interviens, la gouvernance de l'IA est perçue comme un frein. Le raisonnement implicite est toujours le même : gouverner, c'est ralentir, ajouter des validations, des comités, des règles, au moment précis où la concurrence, elle, avance vite. Je démonte ce raisonnement dans ce chapitre, non pas par principe, mais par les chiffres.

Ce que l'adoption sans preuve produit réellement

Si l'opposition vitesse/gouvernance était fondée, on devrait observer que les entreprises qui adoptent le plus vite l'IA générative sont aussi celles qui en tirent le plus de valeur. Ce n'est pas ce que montrent les données disponibles.

Selon le BCG AI Radar (2025), seules 5 % des entreprises dans le monde sont ce que le cabinet appelle « AI-future built », c'est-à-dire qu'elles captent une valeur substantielle de leurs investissements IA. 35 % sont en phase de montée en charge et commencent à en tirer un bénéfice réel. Et **60 % des entreprises ne tirent qu'une valeur minime de l'IA, malgré des investissements significatifs** [1].



Où va la valeur de l'IA en entreprise (BCG, 2025)

McKinsey confirme ce constat sous un autre angle : seulement 39 % des organisations parviennent à attribuer un impact mesurable de l'IA sur leur EBIT, et parmi elles, la majorité rapporte un impact inférieur à 5 %. À peine 6 % des organisations dans le monde sont des « hauts performeurs », avec un impact sur l'EBIT supérieur ou égal à 5 % [2].

Autrement dit : l'adoption est massive (88 % des organisations selon McKinsey, voir Chapitre 1), mais la valeur captée reste rare. La vitesse d'adoption et la vitesse de résultat ne sont pas la même chose. C'est précisément l'écart que la gouvernance de la preuve cherche à combler, non pas en ralentissant l'adoption, mais en la rendant capable de produire un résultat mesurable et défendable.

Le coût, différé mais réel, de l'absence de gouvernance

L'argument le plus solide contre le mythe de la vitesse n'est pas théorique : il est déjà chiffré. Le rapport IBM/Ponemon sur le coût des violations de données (2025) montre que 20 % des

organisations victimes d'une violation en 2025 l'ont été via des usages d'IA non gouvernés (« shadow AI »), avec un surcoût moyen de **670 000 dollars par incident**. 63 % des organisations touchées n'avaient aucune politique de gouvernance IA en place. 97 % manquaient de contrôles d'accès appropriés [3].

C'est le chiffre que je pose systématiquement sur la table en COMEX, et il mérite d'être lu lentement : ce n'est pas le coût de la gouvernance. C'est le coût de son absence.

Sur le terrain juridique, trois décisions rendues en 2025 confirment le même mécanisme à une autre échelle. Un tribunal allemand a condamné OpenAI pour la mémorisation et la reproduction de paroles protégées par le droit d'auteur, rejetant l'exception de fouille de données [4]. Anthropic a accepté un règlement de 1,5 milliard de dollars après avoir constitué une partie de ses données d'entraînement à partir de livres piratés [5]. Dans les deux cas, le problème n'était pas la vitesse d'adoption de l'IA générative par ces entreprises, c'était l'absence de traçabilité sur l'origine et le traitement des données utilisées.

Ces litiges concernent aujourd'hui principalement les fournisseurs de modèles IA (OpenAI, Anthropic, Stability AI), pas encore massivement leurs entreprises clientes. Ma lecture est que le principe juridique qu'ils posent, la mémorisation et la reproduction non tracées engagent la responsabilité, s'appliquera mécaniquement aux entreprises qui déploient des architectures RAG ou des agents IA sur leurs propres contenus et données, sans traçabilité équivalente. C'est un point que le Chapitre 4 développera.

Ce que ce chapitre ne prétend pas démontrer

Je ne prétendrai pas que l'absence de gouvernance ralentit déjà, aujourd'hui, l'ensemble du marché. Ce n'est pas le cas : l'adoption de l'IA générative continue de croître fortement, y compris dans des organisations peu gouvernées. Ce que les données montrent, plus précisément, c'est que :

- la vitesse d'adoption sans gouvernance ne se traduit pas en valeur mesurable pour 60 % des entreprises (BCG) ;
- lorsque l'absence de gouvernance se matérialise en incident, son coût est significatif et documenté (IBM) ;
- les premiers précédents juridiques structurants sont posés (2025), avant même la pleine application du calendrier AI Act (voir Chapitre 5).

Ma thèse n'est donc pas « la gouvernance accélère tout le monde dès aujourd'hui ». Elle est plus précise, et plus défendable : **l'absence de preuve plafonne la valeur captée, et transforme un risque diffus en coût concentré au moment où il se matérialise**. C'est un pari sur la trajectoire, pas une loi déjà vérifiée partout, mais c'est un pari appuyé sur des chiffres, pas sur une conviction.

Conclusion du chapitre

La vraie question n'est pas « faut-il aller vite ou gouverner ». C'est : **peut-on construire, dès maintenant, la preuve qui permet d'aller vite sans être arrêté plus tard ?** Les chapitres suivants (Partie 2, Risque) montrent, à partir de ce que j'observe en mission, où cette preuve manque le plus souvent, et pourquoi ce manque coûte déjà cher à ceux qui l'ignorent.

Sources

- [1] BCG, « Are You Generating Value from AI? The Widening Gap », AI Radar, 30 septembre 2025.
 [2] McKinsey QuantumBlack, « The state of AI in 2025 », novembre 2025.
 [3] IBM (X-Force/Ponemon Institute), « Cost of a Data Breach Report 2025 », 2025.
 [4] Landgericht München I, GEMA c/ OpenAI, affaire n° 42 O 14139/24, 11 novembre 2025.
 [5] Bartz v. Anthropic, règlement transactionnel, août 2025.

Chapitre 3 : La dette algorithmique

PARTIE 2 : RISQUE

La dette technique est un concept que tout DSI connaît : chaque raccourci pris pour livrer plus vite un système d'information se paie, un jour, en temps de refonte, en bugs, en incidents. La dette algorithmique est son équivalent pour l'IA générative, à un détail près, plus grave : elle ne se voit pas dans le code, elle se voit dans les contenus produits, les décisions prises et les données ingérées, sans que personne ne l'ait consignée quelque part.

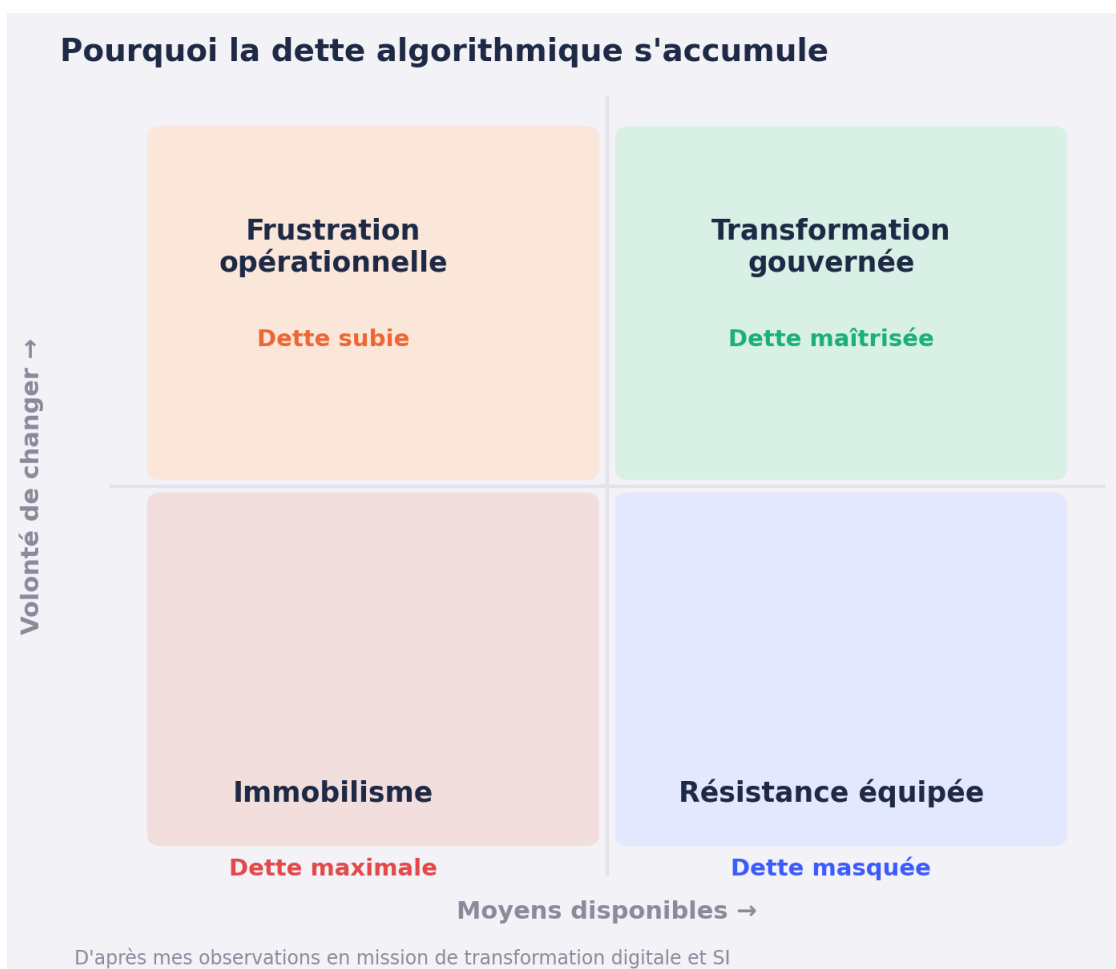
Ce que l'expérience terrain montre, entreprise après entreprise

En quinze ans, dans les organisations que j'ai accompagnées, je n'ai jamais vu la dette algorithmique s'accumuler par accident ni par négligence isolée. Elle s'accumule pour deux raisons récurrentes, presque toujours combinées : **le manque de moyens**, et **le manque de volonté de changement**, que cette réticence vienne des équipes elles-mêmes ou d'autres parties prenantes de l'organisation (direction, autres directions, partenaires).

Ces deux facteurs ne se comportent pas de la même façon, et je constate que les confondre conduit systématiquement à de mauvais diagnostics :

- **Le manque de moyens** (budget, compétences internes, outillage, temps dédié) produit une dette **subie** : l'organisation sait qu'il faudrait gouverner l'usage de l'IA, mais n'en a pas les moyens immédiats. C'est un problème d'arbitrage, qui se résout par la priorisation.
- **Le manque de volonté de changement** produit une dette **choisie**, souvent sans le dire : des équipes qui préfèrent garder leurs habitudes plutôt que d'adopter de nouveaux garde-fous, une direction qui ne veut pas remettre en cause une organisation en silos, un prestataire qui n'a pas intérêt à ce que ses livrables IA soient audités. C'est un problème de gouvernance et de posture managériale, pas de budget.

La combinaison des deux est la plus dangereuse : une organisation qui manque à la fois de moyens et de volonté n'a ni la capacité ni l'intention de traiter sa dette algorithmique, elle continue de l'accumuler jusqu'à ce qu'un incident, un contrôle ou un litige la rende visible de force.

*Pourquoi la dette algorithmique s'accumule*

Cette matrice croise les deux facteurs. Le quadrant que je rencontre le plus souvent n'est pas l'« immobilisme » total (faibles moyens, faible volonté), c'est la « **résistance équipée** » : des organisations qui ont les moyens financiers et techniques de gouverner leur IA, mais où la volonté de changement fait défaut, le plus souvent parce que la gouvernance vient remettre en cause des habitudes de travail bien installées, ou des périmètres de pouvoir entre directions. Le manque de moyens est l'explication la plus souvent avancée en COMEX ; ce n'est pas, dans les organisations que j'accompagne, la cause dominante.

Où se loge concrètement la dette

La dette algorithmique n'est pas un concept abstrait : elle prend des formes précises, que je retrouve d'un audit à l'autre dans le système d'information.

Contenu non tracé. Des visuels, textes ou vidéos générés par IA, publiés sans registre de leur origine, quel outil, quel prompt, quelles sources mobilisées. Le jour où une question de droit d'auteur ou de conformité de marque se pose, personne ne peut répondre.

Décisions automatisées non documentées. Un score de scoring client, une priorisation de leads, un tri de candidatures, produits par un modèle dont personne ne peut expliquer précisément la logique au moment où un client, un candidat ou un régulateur le demande.

Données d'entraînement ou de récupération non gouvernées. Des bases documentaires internes connectées à un outil IA sans revue préalable de ce qu'elles contiennent, données personnelles, informations confidentielles, contenus obsolètes.

Accès non audités. Des connecteurs entre outils IA et systèmes d'entreprise (CRM, ERP, messagerie) mis en place rapidement, sans revue des droits accordés, un sujet développé en détail au Chapitre 4.

Le coût, déjà visible

Ce n'est plus un risque théorique, et ce n'est plus un argument que j'ai besoin de défendre longtemps devant un COMEX. Le rapport IBM/Ponemon 2025 chiffre à 670 000 dollars le surcoût moyen d'une violation de données liée au shadow AI, avec 97 % des organisations touchées ne disposant pas de contrôles d'accès appropriés [1]. Les décisions GEMA c/ OpenAI et Bartz v. Anthropic (Chapitre 2) montrent que l'absence de traçabilité sur l'origine des contenus et données peut coûter jusqu'à 1,5 milliard de dollars lorsqu'elle explose au tribunal [2]. Le précédent Moffatt v. Air Canada rappelle qu'une entreprise reste responsable de ce que dit son propre système IA, même quand elle ne l'a pas explicitement validé mot pour mot [3].

Dans tous ces cas, la cause profonde est la même : une dette accumulée par manque de moyens ou de volonté, restée invisible jusqu'à ce qu'elle devienne un incident, un litige ou une sanction.

Conclusion du chapitre

La dette algorithmique n'est ni une fatalité technique ni un problème réservé aux grandes entreprises exposées médiatiquement. C'est le résultat cumulé de choix organisationnels, arbitrages de moyens, résistances au changement, que la plupart des comités de direction n'ont jamais formalisés comme tels. Le premier acte de gouvernance que je recommande n'est donc pas technique : c'est de nommer cette dette, de savoir dans quel quadrant se situe votre organisation, et d'admettre laquelle des deux causes, moyens ou volonté, pèse le plus. Le Chapitre 4 descend d'un cran, au niveau de l'architecture technique, pour montrer précisément où cette dette se loge dans les systèmes RAG et les agents IA connectés via MCP.

Sources

[1] IBM (X-Force/Ponemon Institute), « Cost of a Data Breach Report 2025 », 2025.

[2] Landgericht München I, GEMA c/ OpenAI, 11 novembre 2025 ; Bartz v. Anthropic, règlement août 2025.

[3] Civil Resolution Tribunal of British Columbia, Moffatt v. Air Canada, 14 février 2024.

Chapitre 4 : RAG et MCP : les nouveaux angles morts

PARTIE 2 : RISQUE

Le Chapitre 3 a montré que la dette algorithmique se loge dans des choix organisationnels, moyens et volonté. Je descends ici au niveau de l'architecture technique, là où cette dette prend une forme précise et où, à ce jour, très peu des organisations que j'accompagne ont posé un cadre de gouvernance : les systèmes RAG (Retrieval-Augmented Generation) et les agents IA connectés via des protocoles comme MCP (Model Context Protocol).

C'est le chapitre le plus technique de ce livre blanc, et il l'est délibérément : la gouvernance de l'IA ne se joue pas seulement au niveau des principes et des chartes, elle se joue dans l'architecture, et c'est précisément là que la plupart des discours sur « l'éthique de l'IA » s'arrêtent avant d'avoir commencé.

Le RAG : rendre l'IA pertinente, au prix de nouveaux risques

Un système RAG permet à un modèle d'IA générative de s'appuyer sur des documents internes, contrats, bases de connaissances, CRM, documentation produit, pour générer des réponses ancrées dans les données réelles de l'entreprise, plutôt que dans les seules connaissances générales du modèle. C'est l'architecture la plus répandue pour rendre l'IA générative utile en entreprise.

Elle introduit cependant des points de rupture de la preuve, et ce sont les quatre questions que je pose en premier lors d'un audit :

- **Traçabilité de la source** : quand une réponse est générée, peut-on remonter exactement au document qui l'a alimentée ?
- **Droits d'accès** : le système qui indexe les documents respecte-t-il les mêmes droits d'accès que les utilisateurs humains, ou tout le monde voit-il, via l'IA, des documents auxquels il n'aurait normalement pas accès ?
- **Fraîcheur des données** : le corpus indexé est-il à jour, ou l'IA répond-elle avec des informations obsolètes sans que l'utilisateur le sache ?
- **Intégrité du corpus** : les documents utilisés comme source ont-ils pu être altérés ou injectés de façon malveillante ?

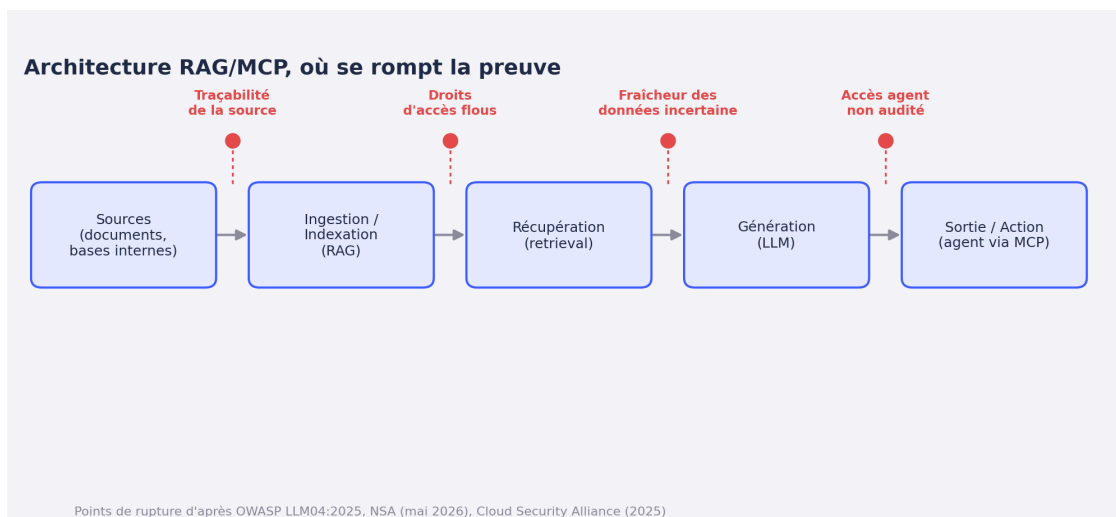
Ce dernier point n'est pas hypothétique. L'OWASP, référentiel de sécurité applicative reconnu, classe l'empoisonnement de données/modèle comme le quatrième risque majeur des systèmes LLM en 2025, identifiant explicitement les sources externes non vérifiées, dont les corpus RAG, comme vecteur à haut risque [1]. Des travaux de recherche académique publiés en 2025 (dont une communication à l'ACM Web Conference) démontrent qu'un seul document empoisonné suffit, dans certains cas, à influencer significativement les réponses générées, un phénomène documenté sous le nom de « one-shot dominance » [2].

Le MCP : quand l'IA agit, pas seulement répond

Le Model Context Protocol et les protocoles équivalents permettent à des agents IA de se connecter directement à des systèmes d'entreprise, CRM, ERP, messagerie, outils internes, pour y lire des informations, mais aussi pour y agir : créer un enregistrement, envoyer un message, déclencher un processus. C'est le saut qualitatif de 2025-2026 : l'IA générative cesse d'être un outil consulté pour devenir un acteur du système d'information.

Ce saut change la nature du risque. Un agent IA connecté à un CRM ou un ERP sans cloisonnement explicite des droits est, du point de vue de la sécurité, équivalent à un compte administrateur non audité.

La NSA (agence de sécurité américaine), dans un document technique publié en mai 2026, qualifie la conception du protocole MCP de « flexible et sous-spécifiée », identifiant des risques structurels : exécution de code arbitraire via des entrées non contraintes, authentification optionnelle sans contrôle d'accès basé sur les rôles standardisé au niveau du protocole, absence de gestion de cycle de vie des jetons d'accès [3]. Des cas réels d'exploitation ont déjà été documentés : exfiltration de données via un serveur MCP malveillant, collision de noms d'outils permettant le chargement de code arbitraire, et au moins une vulnérabilité critique (CVE-2025-49596) de type exécution de code à distance dans des outils de développement connectés via MCP [3]. La Cloud Security Alliance recense six familles de risques équivalents, dont les attaques dites « rug-pull », un serveur MCP approuvé initialement qui modifie ensuite son comportement sans repasser par une validation [4].



Architecture RAG/MCP, où se rompt la preuve

Ce que ce constat implique, et ce qu'il n'implique pas

Je veux être précis sur la portée de ces éléments : les cas documentés à ce jour sont majoritairement des démonstrations de faisabilité en environnement de recherche ou de sécurité, pas un décompte massif d'incidents réels subis en production par des entreprises. Le corpus est cohérent avec un sujet encore récent, probablement pas encore couvert par une base de données d'incidents publics comparable à celles qui existent pour d'autres types de cyberattaques. Je ne prétends donc pas qu'une vague d'incidents RAG/MCP est déjà en cours. J'affirme quelque chose de plus précis et de vérifiable : **le risque est structurellement documenté par des institutions de sécurité de premier plan (NSA, OWASP, Cloud Security Alliance), avant même que les statistiques d'incidence ne soient disponibles**, ce qui signifie qu'agir maintenant est une anticipation raisonnable, pas une réaction à une crise déjà là.

C'est précisément la fenêtre où la gouvernance a le plus de valeur : avant que le risque documenté devienne un incident chiffré.

Ce que la gouvernance de la preuve exige, concrètement

J'en retiens trois principes, directement dérivés des sources ci-dessus, qui doivent structurer toute architecture RAG/MCP en entreprise (développés opérationnellement au Chapitre 7) :

1. **Traçabilité par construction** : chaque réponse générée doit pouvoir être reliée à ses sources ; chaque action d'un agent doit être journalisée.
2. **Accès restreint par défaut** : un système RAG hérite des droits d'accès existants plutôt que de les contourner ; un agent MCP démarre avec le minimum de privilèges nécessaires, étendus seulement avec preuve de gouvernance.

3. **Revalidation périodique** : les sources indexées, les outils connectés et les droits accordés sont revus régulièrement, pas validés une fois puis oubliés.

Conclusion du chapitre

La gouvernance de l'IA générative ne se limite pas à une charte d'usage ou à une politique de conformité. Elle se joue, concrètement, dans la façon dont un système RAG est construit et dont un agent MCP est cloisonné. C'est le niveau de granularité auquel la plupart des discours généralistes sur « l'IA éthique » ne descendent jamais, et c'est précisément là que se joue, en 2026, la différence entre une organisation qui maîtrise son IA et une organisation qui la subit.

Sources

[1] OWASP GenAI Security Project, « LLM04:2025 Data and Model Poisoning », 5 mai 2025.

[2] ACM Web Conference 2025, « Traceback of Poisoning Attacks to Retrieval-Augmented Generation » ; arXiv 2505.11548, « One Shot Dominance », 2025.

[3] NSA, « Security Design Considerations for AI-Driven Automation (MCP Security) », Cybersecurity Information Sheet, mai 2026.

[4] Cloud Security Alliance, « Agentic MCP Security Best Practices Guide », 2025.

Chapitre 5 : Le mur réglementaire

PARTIE 2 : RISQUE

Il existe une façon paresseuse de parler de l'AI Act dans un livre blanc : citer sa date d'entrée en vigueur, agiter la menace de sanctions, et passer au chapitre suivant. Je prends ici le chemin inverse : j'explique précisément ce qui change, quand, et surtout ce qui n'est **pas encore** arrivé, parce que se tromper sur ce point coûterait cher à la crédibilité de ce document face à un public de DSI et de juristes.

Un calendrier qui vient d'être révisé

Le règlement européen sur l'IA (2024/1689) est entré en application par étapes depuis février 2025. Mais son calendrier a été substantiellement modifié à la mi-2026, via un texte dit « Digital Omnibus » adopté par le Conseil de l'Union européenne le 29 juin 2026 [1]. J'y vois un point de vigilance majeur : toute publication antérieure à cette date citant « août 2026 » comme échéance pour les obligations relatives aux systèmes à haut risque est aujourd'hui obsolète.



Calendrier AI Act, révisé mi-2026

Le calendrier à jour se lit ainsi : les interdictions de base et les obligations de littératie IA sont applicables depuis février 2025 ; la gouvernance des modèles à usage général (GPAI) depuis août 2025 ; l'application générale du règlement a débuté en août 2026. Les obligations les plus lourdes, pour les systèmes à haut risque autonomes, ont été **reportées d'un an**, du 2 août 2026 initialement prévu au 2 décembre 2027. Les systèmes à haut risque intégrés à des produits déjà réglementés (dispositifs médicaux, ascenseurs) suivront en août 2028 [1].

Ce report ne change rien à la logique du calendrier : il donne simplement un an de plus aux organisations pour se préparer. C'est une fenêtre à saisir, pas une raison de repousser le sujet.

Ce que le règlement exige concrètement

L'approche de l'AI Act est fondée sur les risques : plus un système IA peut avoir un impact significatif sur les personnes, plus les obligations sont lourdes. Pour les systèmes à haut risque, cela se traduit par une documentation technique, une tenue de registres permettant la traçabilité, et une information claire des déployeurs sur les capacités et limites du système (article 13). Pour l'IA générative spécifiquement, une obligation distincte impose d'informer les utilisateurs qu'ils interagissent avec une IA et de marquer le contenu synthétique (article 50) [2].

Je résume ici l'esprit de ces articles ; pour toute citation littérale opposable, référez-vous au texte officiel du règlement.

Ce qui n'est pas encore arrivé, et pourquoi ça ne doit pas rassurer

Il serait tentant, pour un COMEX pressé, de conclure que le sujet peut attendre puisque les obligations les plus lourdes sont reportées à fin 2027. Deux faits contredisent cette lecture, et je les mets systématiquement sur la table en mission.

D'abord, un précédent sanctionnateur existe déjà en Europe : l'autorité italienne de protection des données a infligé, en janvier 2025, une amende de 15 millions d'euros à OpenAI pour violation du RGPD, absence de base légale adéquate, absence de vérification d'âge, manquement à la notification de violation [3]. Ce n'est pas un cas isolé au sens strict : c'est le premier d'une série qui s'ouvre, sous un cadre RGPD déjà pleinement applicable, indépendamment du calendrier spécifique à l'AI Act.

Ensuite, côté France, la CNIL n'a **pas encore** prononcé de sanction visant spécifiquement un usage d'IA générative en entreprise, mais elle a publié, dès février et juillet 2025, un corpus complet de recommandations opérationnelles sur l'IA et le RGPD (finalité, minimisation, information des personnes, mémorisation des modèles, sécurité des environnements de développement) [4]. Autrement dit : le référentiel d'attente existe déjà, même en l'absence de contentieux. Une organisation qui attendrait une sanction française pour se mettre en conformité se base sur une absence de précédent, pas sur une absence de règle.

Pourquoi la conformité anticipée est un avantage, pas une contrainte

La thèse que je défends dans ce livre blanc (voir Ch.2) est que la mise en conformité intégrée dès la conception du SI coûte moins cher et va plus vite que la mise en conformité a posteriori. L'AI Act en offre une démonstration directe : les organisations qui construisent, dès maintenant, la traçabilité et la documentation exigées par les articles 13 et 50 n'auront rien à refaire lorsque les obligations haut risque entreront en application fin 2027. Celles qui attendent devront, dans un délai contraint, documenter rétroactivement des systèmes déjà en production, un exercice nettement plus coûteux qu'une conception initiale gouvernée.

Conclusion du chapitre

Le mur réglementaire n'est ni imminent au sens d'une sanction immédiate, ni lointain au sens où il pourrait être ignoré. Il est engagé, documenté, daté avec précision, et le report de calendrier de 2026 doit être lu comme un délai de préparation, pas comme un signal de moindre urgence. C'est le dernier chapitre de la Partie RISQUE : dans les Chapitres 6, 7 et 8 qui suivent (Partie TRANSFORMATION), je montre comment transformer cette contrainte en organisation gouvernée, plutôt que de la subir.

Sources

[1] Conseil de l'Union européenne, communiqué « AI: Council gives final green light to simplify and streamline rules », 29 juin 2026.

[2] Règlement (UE) 2024/1689, articles 13 et 50 (texte annoté).

[3] Garante per la protezione dei dati personali, sanction OpenAI, janvier 2025.

[4] CNIL, recommandations IA et RGPD, 7 février 2025 et 22 juillet 2025.

Chapitre 6 : Sortir des silos

PARTIE 3 : TRANSFORMATION

Les cinq premiers chapitres ont posé un constat (l'IA est partout, la vitesse sans preuve ne paie pas) et un risque (la dette algorithmique, ses angles morts techniques, le mur réglementaire). J'ouvre ici la Partie Transformation par une question organisationnelle, trop souvent éludée : pourquoi la gouvernance de l'IA échoue-t-elle, la plupart du temps, alors que les principes qu'elle porte ne sont pas contestés sur le fond ?

La réponse tient en un mot : les silos.

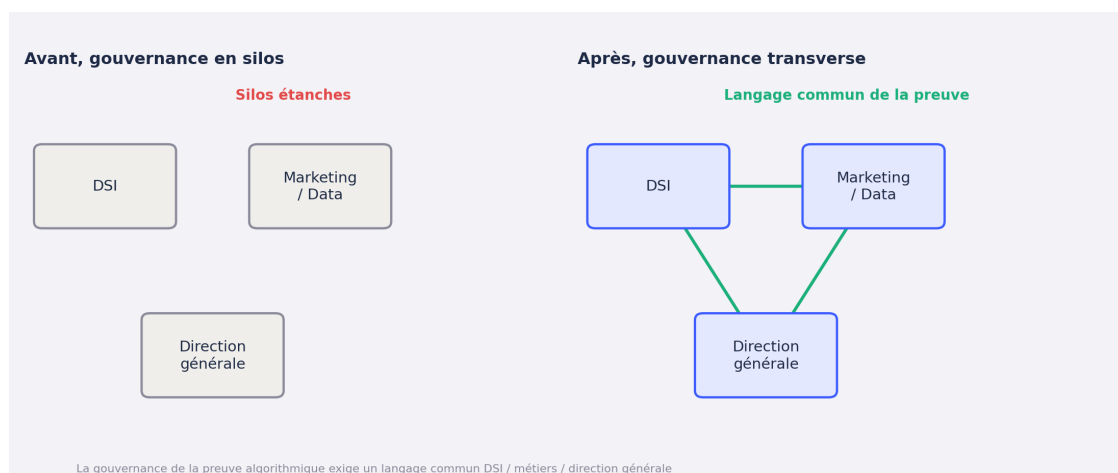
Le renvoi mutuel de responsabilité

Dans la majorité des organisations que j'accompagne, la gouvernance de l'IA générative est un sujet que chaque fonction renvoie à une autre. La DSI considère que la gouvernance des usages relève des métiers qui déploient l'IA. Le marketing ou les métiers considèrent que la gouvernance technique, accès, traçabilité, sécurité, relève de la DSI. Le juridique attend d'être sollicité plutôt que d'imposer un cadre en amont. Le résultat : un sujet transverse par nature, traité par personne dans son ensemble.

Ce constat rejoint directement ce que j'ai décrit au Chapitre 3 sur les causes de la dette algorithmique : le manque de volonté de changement, qu'il vienne des équipes ou d'autres parties prenantes, est souvent plus déterminant que le manque de moyens. Les silos organisationnels sont l'un des mécanismes principaux par lesquels ce manque de volonté se traduit concrètement : chaque fonction préserve son périmètre plutôt que de construire un langage commun, et la dette s'accumule dans les interstices que ce cloisonnement laisse ouverts.

Ce que le cadre réglementaire impose déjà comme langage commun

Le Chapitre 5 a montré que le RGPD et l'AI Act imposent déjà des obligations qui, par construction, ne peuvent pas être portées par une seule fonction : la documentation technique et la traçabilité (article 13 de l'AI Act) supposent une collaboration DSI-métiers ; l'information des personnes concernées et la gestion des droits (RGPD, recommandations CNIL) supposent une collaboration juridique-métiers ; la sécurité des environnements de développement IA (recommandations CNIL de juillet 2025) suppose une collaboration DSI-sécurité. Le cadre normatif ne laisse déjà plus la place à une gouvernance en silo, mais les organisations, elles, n'ont pas toujours suivi.



Avant / après, gouvernance en silos vs gouvernance transverse

Pourquoi ce pont est plus facile à construire qu'il n'y paraît

L'exemple le plus parlant, dans mon expérience, ne concerne pas encore l'IA : il lui est antérieur, et c'est justement ce qui en fait une bonne démonstration. Je suis arrivé dans plusieurs organisations où le point de départ était le même : des fichiers Excel « volants », 50 000 lignes, parfois vieux de quinze ans, que les équipes se passaient à la main d'un service à l'autre pour faire vivre l'information. Traiter une donnée nouvelle relevait du parcours du combattant : personne n'avait de vue d'ensemble, chaque service ne voyait que sa version du fichier, à un instant donné, sans savoir si elle était à jour.

Le premier chantier, dans ces situations, n'a jamais été un chantier IA. C'était la mise en place d'un **ERP avec une passerelle vers l'ensemble des services**, pour réduire le coût de circulation de l'information, moins de ressaisie, moins de versions divergentes, une source de vérité unique. C'est très exactement le même problème que pose l'IA générative aujourd'hui, à une différence près : le fichier Excel qui circule sans traçabilité est devenu un prompt, une base RAG ou un agent connecté qui circule sans traçabilité. Le symptôme a changé de forme, la cause organisationnelle est identique, l'absence de langage et d'infrastructure communs entre services.

C'est cette continuité, de l'Excel non gouverné hier à l'IA non gouvernée aujourd'hui, qui me sert le plus souvent à convaincre un comité de direction déjà familier du problème ERP/data des années 2010-2020 : la gouvernance de l'IA n'est pas un sujet nouveau et exotique, c'est la suite d'un problème qu'ils ont déjà vécu, et déjà résolu une fois, avec les mêmes leviers (source de vérité unique, passerelle inter-services, réduction du coût de circulation de l'information).

Le vocabulaire commun que je propose dans ce livre blanc est celui de la **preuve** : quelle que soit la fonction, la question « peut-on prouver ce que ce système IA a fait ? » se pose dans les mêmes termes pour un DSI, un directeur marketing, un juriste ou un DG. C'est un langage qui ne privilégie aucun silo, il les traverse tous, et l'exemple ERP ci-dessus montre qu'un langage équivalent a déjà fait ses preuves sur un problème structurellement identique.

Ce qu'une gouvernance transverse change concrètement

Une organisation qui sort des silos ne crée pas nécessairement un nouveau poste ou une nouvelle direction (ce sujet, la répartition des responsabilités, est traité en détail au Chapitre 9). Elle change d'abord trois pratiques :

- **Un registre unique des usages IA**, alimenté par toutes les fonctions, plutôt que des inventaires séparés (ou inexistants) par direction.
- **Un point de décision commun** pour toute nouvelle intégration IA touchant à des données ou des systèmes partagés, pas un veto centralisé, mais un passage obligé documenté.
- **Un vocabulaire partagé** (traçabilité, accès, sourcing, réversibilité) utilisé de la même façon par la DSI, les métiers et la direction générale.

Conclusion du chapitre

Le problème de la gouvernance IA n'est pas d'abord technique : il est organisationnel, et il se loge précisément là où les Chapitres 3 et 4 ont identifié la dette, dans les manques de moyens, mais surtout dans les manques de volonté de changement qui traversent les silos. Dans le chapitre suivant, je montre comment ce pont s'incarne concrètement, à travers une brique que je pratique directement : l'intégration opérationnelle de l'IA dans les process métier.

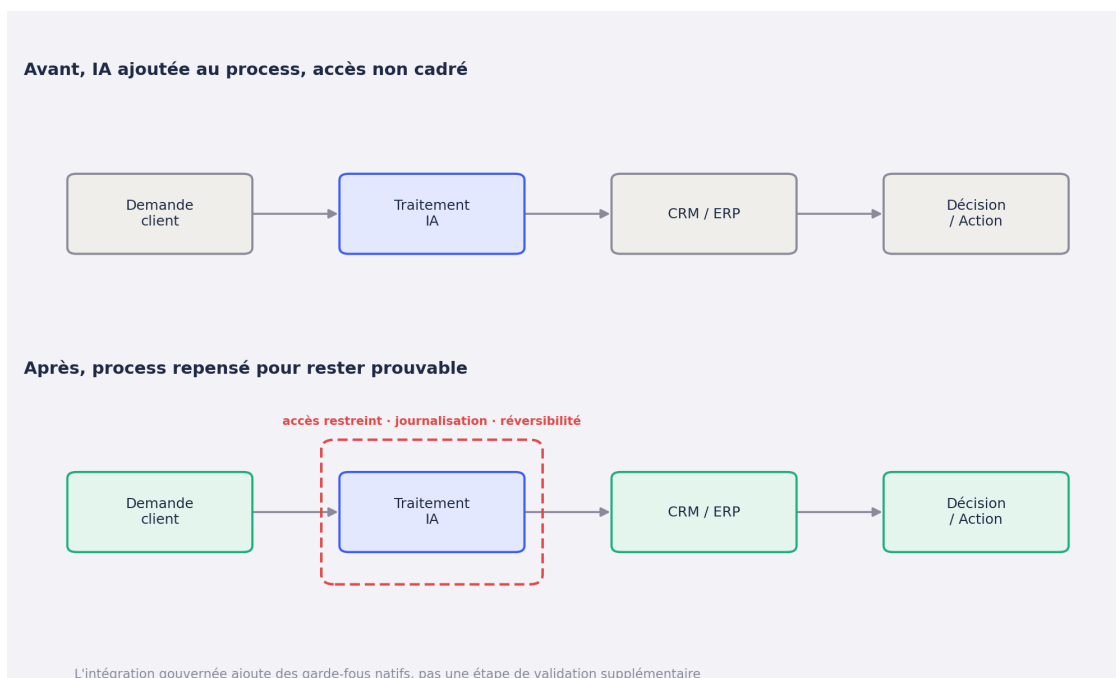
Chapitre 7 : Intégrer l'IA dans les nouveaux process sans perdre la maîtrise

PARTIE 3 : TRANSFORMATION

Je clôt la Partie Transformation en descendant du principe à la pratique : comment déployer concrètement RAG, agents IA et automatisations dans les process d'entreprise, CRM, ERP, service client, production de contenu, sans reproduire, à l'échelle de l'IA générative, l'erreur que l'exemple ERP du Chapitre 6 a montré à l'échelle de la donnée.

L'erreur la plus fréquente : ajouter l'IA au process existant

La tentation la plus commune, quand une organisation déploie l'IA générative dans un process métier, est de la traiter comme une brique supplémentaire : le process reste identique, on y insère un point de traitement IA, et on avance. Cette approche fonctionne à court terme, elle produit rapidement des gains visibles, mais elle reproduit exactement le schéma qui a créé la dette algorithmique décrite au Chapitre 3 : l'IA s'infiltre dans un process qui n'a pas été pensé pour rester prouvable une fois qu'elle y est intégrée.



Avant/après, intégration de l'IA dans un process métier

La différence entre les deux schémas n'est pas la présence de garde-fous en plus : c'est que, dans le second, ces garde-fous sont **natifs**, pas ajoutés après coup. Un accès restreint par défaut, une journalisation systématique et une réversibilité prévue dès la conception ne ralentissent pas le process, ils évitent d'avoir à le refaire quand la preuve sera exigée (par un client, un régulateur, un tribunal).

Trois principes d'architecture, dérivés du Chapitre 4

Le Chapitre 4 a établi trois principes pour les architectures RAG/MCP au niveau technique. Je les applique ici au niveau du process métier :

1. Accès restreint par défaut. Un agent IA connecté à un CRM ou un ERP ne doit hériter que des droits strictement nécessaires à sa tâche, pas des droits de la personne ou du service qui l'a configuré. C'est le principe de moindre privilège, déjà appliqué en cybersécurité, mais rarement transposé aux intégrations IA.

2. Traçabilité native. Chaque action d'un agent IA sur un système d'entreprise, création d'un enregistrement, envoi d'un message, modification d'une donnée, doit être journalisée avec la même rigueur qu'une action humaine. Un agent qui « agit au nom de » l'entreprise engage sa responsabilité exactement comme un salarié (voir *Moffatt v. Air Canada*, Chapitre 3), l'absence de journal d'action rend cette responsabilité impossible à documenter en cas de litige.

3. Réversibilité prévue. Toute automatisation IA d'un process doit prévoir, dès sa conception, comment revenir en arrière si le résultat est erroné ou contesté, pas comme un correctif d'urgence, mais comme une fonctionnalité du process lui-même.

Ce que cela change dans un projet CRM/ERP concret

Le Chapitre 6 l'a montré avec l'exemple des fichiers Excel remplacés par un ERP doté d'une passerelle inter-services : le premier chantier de fond n'est presque jamais l'IA elle-même, c'est l'architecture qui la rend gouvernable. Le même principe s'applique à l'intégration de RAG et d'agents IA dans un CRM/ERP existant : avant de connecter un agent IA à un système d'information, la vraie question n'est pas « quel outil choisir » mais « ce système peut-il aujourd'hui produire une preuve de ce qui s'y passe », la même question que pose une passerelle inter-services bien conçue.

Buy vs build : une fausse question prioritaire

Beaucoup d'organisations abordent l'intégration de l'IA dans leurs process par la question de l'outil : SaaS packagé ou développement interne. Cette question est secondaire face à celle de la preuve : un outil SaaS d'IA sans garanties contractuelles de traçabilité et de réversibilité est plus risqué qu'un développement interne moins sophistiqué mais entièrement auditable. Le choix technologique doit suivre l'exigence de preuve, pas l'inverse.

Conclusion du chapitre

Intégrer l'IA dans les process métier n'est pas un projet d'outillage, c'est un projet d'architecture, la même discipline que celle qui a transformé des fichiers Excel non gouvernés en ERP structuré, appliquée cette fois à des systèmes qui agissent, pas seulement qui stockent. Dans la Partie Gouvernance qui suit (Chapitres 8-9), je formalise ces principes en un cadre complet et en une répartition claire des responsabilités.

Chapitre 8 : Le cadre de la preuve algorithmique

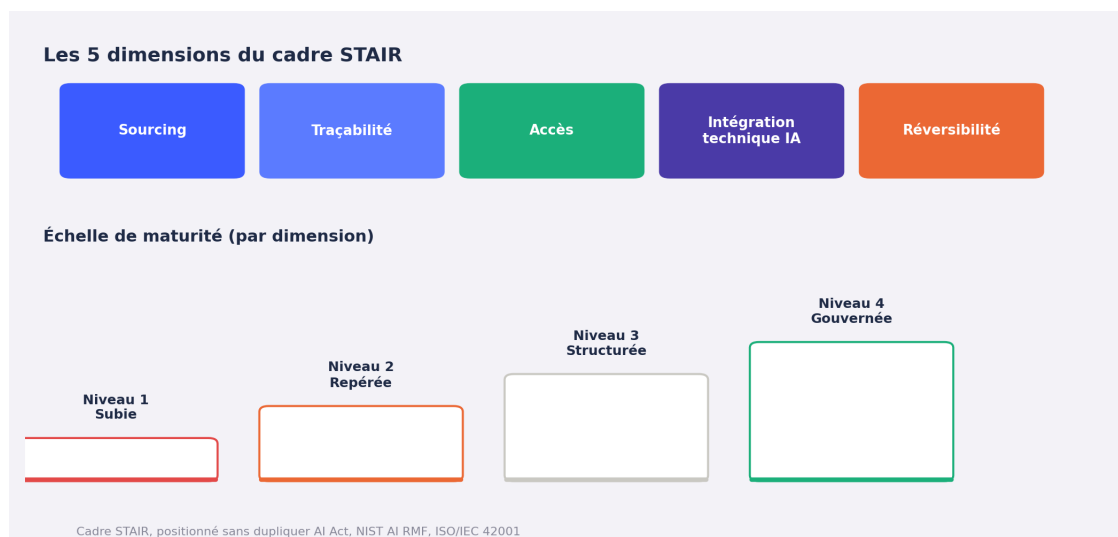
PARTIE 4 : GOUVERNANCE

Les sept chapitres précédents ont construit, un par un, les éléments d'un même constat : l'IA est diffuse (Ch.1), sa vitesse d'adoption ne suffit pas (Ch.2), son absence de gouvernance produit une dette mesurable (Ch.3), cette dette se loge précisément dans des architectures techniques identifiables (Ch.4), un calendrier réglementaire formalise déjà l'exigence de preuve (Ch.5), le problème est avant tout organisationnel (Ch.6), et l'intégration de l'IA dans les process métier (Ch.7) montre que la gouvernance de la preuve est déjà praticable, pas seulement théorique.

Ce chapitre assemble ces éléments en un cadre structuré : les dimensions à gouverner, et les niveaux de maturité qui permettent à une organisation de se situer.

Cinq dimensions, pas dix principes

La plupart des cadres de gouvernance IA existants (voir plus bas) multiplient les principes au point de devenir inopérants pour une équipe opérationnelle. J'en retiens cinq, choisies parce que chacune correspond à un chapitre développé précédemment et à une preuve concrète déjà établie :



Les 5 dimensions du cadre de la preuve algorithmique

- **Sourcing**, les données et contenus mobilisés sont-ils vérifiés, à jour et légitimement utilisables ? (Ch.4, Ch.5, droit d'auteur)
- **Traçabilité**, peut-on relier tout contenu, décision ou action IA à son origine ? (Ch.4, art. 13 AI Act)
- **Accès**, les droits accordés aux systèmes et agents IA sont-ils restreints par défaut et documentés ? (Ch.4, Ch.7)
- **Intégration technique IA**, l'architecture (RAG, MCP, agents) est-elle conçue pour rester gouvernable dans la durée, pas seulement fonctionnelle au lancement ? (Ch.4, Ch.7)
- **Réversibilité**, une décision ou une action IA peut-elle être annulée ou corrigée sans reconstruire le process ? (Ch.7)

Cet ordre n'est pas arbitraire : les initiales des cinq dimensions forment l'acronyme **STAIR**, l'escalier en anglais, et c'est exactement l'image que je retiens en mission. On ne gouverne pas l'IA d'un bloc, on monte les marches une par une, dimension par dimension, à travers les quatre niveaux de maturité décrits ci-dessous.

Chaque dimension se décline sur une échelle de maturité à quatre niveaux, applicable dimension par dimension, une organisation peut être « Structurée » sur la traçabilité et « Subie » sur l'intégration technique, et c'est en réalité la situation que je rencontre le plus souvent dans les organisations que j'accompagne (voir la matrice moyens/volonté, Ch.3).

Niveau 1, Subie. Aucune preuve n'est construite ; la dette s'accumule silencieusement. C'est l'état par défaut d'une organisation qui n'a pas encore nommé le problème.

Niveau 2, Repérée. La dette est cartographiée (un registre des usages IA existe, au moins partiellement) mais pas encore traitée de façon systématique.

Niveau 3, Structurée. Des procédures et une architecture existent pour chaque dimension, accès restreints, journalisation, sourcing documenté, mais leur application reste inégale selon les équipes ou les systèmes.

Niveau 4, Gouvernée. La preuve est continue, auditable à tout moment, et devient un argument de confiance vis-à-vis des clients, partenaires et régulateurs, l'avantage concurrentiel évoqué au Chapitre 11.

Se positionner sans dupliquer l'existant

Ce cadre ne remplace pas les référentiels déjà établis, il les rend actionnables à l'échelle d'une organisation qui n'a pas les moyens de déployer un programme de conformité complet. L'AI Act (articles 13 et 50) impose des obligations réglementaires pour les systèmes à haut risque ; le NIST AI Risk Management Framework structure la gestion des risques IA en quatre fonctions (Govern, Map, Measure, Manage) ; ISO/IEC 42001 formalise un système de management de l'IA selon une logique d'amélioration continue (PDCA) [1].

Le cadre STAIR se positionne en dessous de ces référentiels, comme une **grille de diagnostic opérationnel** : les cinq dimensions et les quatre niveaux de maturité permettent à une équipe dirigeante de savoir, en une conversation, où elle se situe, avant de mobiliser, si nécessaire, les référentiels plus lourds pour la mise en conformité formelle.

Ce que ce chapitre prépare

J'utilise ce cadre en mission sous la forme d'une grille d'auto-évaluation, une ligne par dimension et par niveau de maturité, que je remplis avec l'équipe dirigeante en une séance de travail. Ce chapitre en pose les fondations conceptuelles ; le Chapitre 9 complète la dimension organisationnelle en répondant à la question qui reste ouverte : qui, dans l'entreprise, porte chacune de ces cinq dimensions ?

Sources

[1] NIST AI Risk Management Framework (AI RMF 1.0, 2023) et profil génératif (AI 600-1, 2024) ; ISO/IEC 42001:2023.

Chapitre 9 : Qui gouverne quoi

PARTIE 4 : GOUVERNANCE

Un cadre de gouvernance sans répartition claire des responsabilités reste un document, pas une pratique. Ce chapitre reforge la Partie Gouvernance en répondant à la question la plus concrètement bloquante identifiée depuis le Chapitre 6 : dans une organisation réelle, qui porte chacune des cinq dimensions du cadre présenté au Chapitre 8 ?

Le piège du renvoi mutuel, une dernière fois

Le Chapitre 6 a montré que le renvoi mutuel de responsabilité entre DSI, métiers et juridique est l'un des principaux mécanismes par lesquels la dette algorithmique s'accumule. Une matrice de responsabilité n'élimine pas ce renvoi par magie, mais elle le rend visible, ce qui est la première étape pour le corriger. Quand une organisation ne parvient pas à remplir cette matrice pour son propre usage de l'IA, elle a, de fait, déjà répondu à la question : personne n'y est aujourd'hui responsable de la preuve algorithmique.

Qui gouverne quoi, matrice RACI simplifiée				
	DG	DSI	Métiers	Juridique / Conformité
Sourcing	I	R	A	C
Traçabilité	I	R	C	A
Accès	A	R	C	C
Intégration technique IA	I	R/A	C	I
Réversibilité	A	R	C	I

R = Responsable de l'exécution A = Décisionnaire final C = Consulté I = Informé

Qui gouverne quoi, matrice RACI simplifiée

Lecture de la matrice

Cette matrice n'est pas une réponse universelle, elle doit être adaptée à la structure réelle de chaque organisation, mais elle pose une logique de répartition qui évite les deux écueils que je rencontre le plus souvent : la centralisation excessive (tout remonte à la DSI ou à la direction générale, ce qui bloque l'adoption) et la dilution excessive (chaque métier gère sa propre gouvernance, ce qui recrée les silos du Chapitre 6).

Quatre principes de lecture :

- **La DSI est "Responsable de l'exécution" (R) sur les cinq dimensions**, elle construit et opère l'infrastructure. Sur l'intégration technique IA, et sur elle seule, elle est à la fois exécutante et décisionnaire (R/A) : c'est le seul domaine où le choix d'architecture lui revient en propre.
- **Les métiers sont "Décisionnaires" (A) sur le sourcing**, ce sont eux qui savent quelles données et quels contenus sont légitimes à mobiliser dans leur domaine.
- **La direction générale reste "Accountable" sur l'accès et la réversibilité**, parce que ce sont les deux dimensions où un mauvais arbitrage engage directement la responsabilité de l'entreprise, pas seulement sa performance opérationnelle.

- **Le juridique et la conformité sont "Décisionnaires" (A) sur la traçabilité**, parce que c'est la dimension qui produit la preuve opposable le jour d'un contrôle ou d'un litige. La DSI en construit les journaux, mais c'est au juridique de dire ce qui doit être conservé, sous quelle forme et pendant combien de temps.

Ce que cette répartition change concrètement

L'exemple que j'ai décrit au Chapitre 6 (fichiers Excel remplacés par un ERP avec passerelle inter-services) montre précisément ce que produit l'absence de répartition claire : quand personne n'est identifié comme responsable de la circulation de l'information entre services, chacun gère sa propre version, et la dette s'accumule invisiblement jusqu'à ce qu'un chantier structurant (l'ERP, hier ; la gouvernance IA, aujourd'hui) devienne nécessaire dans l'urgence plutôt que dans l'anticipation.

Cette répartition correspond à une structure standard. Dans certaines organisations, un rôle de coordination transverse explicite fonctionne mieux qu'une matrice DSI/métiers classique, et c'est un arbitrage à faire au cas par cas.

Ce qu'il ne faut pas faire : créer un poste avant de clarifier les responsabilités

L'erreur que je vois le plus souvent consiste à répondre au problème de gouvernance par la création d'un poste (« Chief AI Officer », « responsable IA ») avant même d'avoir clarifié qui décide de quoi. Un poste sans matrice de responsabilité claire en amont devient soit un goulot d'étranglement (tout remonte à cette personne), soit un poste sans autorité réelle (les décisions continuent de se prendre ailleurs). La matrice doit précéder l'organigramme, pas l'inverse.

Conclusion de la Partie Gouvernance

Les Chapitres 8 et 9 ferment la boucle ouverte au Chapitre 1 : l'IA a quitté le pilote, elle est partout dans le système d'information, et la question n'est plus de savoir si elle doit être gouvernée mais comment. Le cadre à cinq dimensions (Ch.8) répond à la question du **quoi** ; la matrice de responsabilité (Ch.9) répond à la question du **qui**. Il reste une dernière question, la plus attendue par un COMEX : **par où commencer, concrètement, dès demain matin**. C'est l'objet de la Partie 5, Action, qui clôt ce livre blanc.

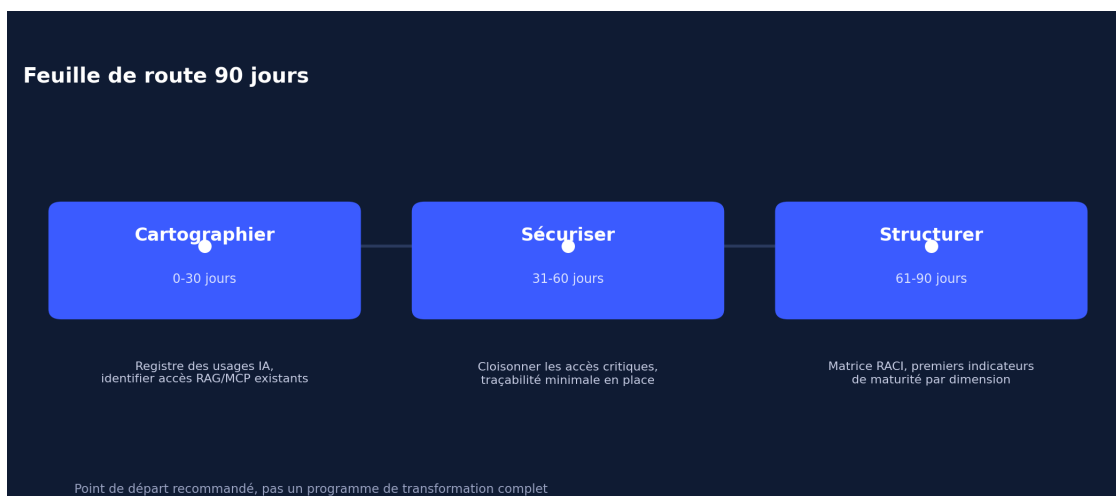
Chapitre 10 : La feuille de route 90 jours pour un COMEX

PARTIE 5 : ACTION

Neuf chapitres ont posé un constat, un risque, une transformation et un cadre de gouvernance. Je réponds ici à la question la plus attendue d'un COMEX à ce stade : par où commencer, concrètement, sans attendre un programme de transformation pluriannuel.

Pourquoi 90 jours, pas 12 mois

Le Chapitre 3 a montré que la dette algorithmique s'accumule pour deux raisons : le manque de moyens et le manque de volonté de changement. Un programme de transformation à 12 ou 18 mois donne du temps aux deux obstacles pour se réinstaller. Une feuille de route à 90 jours, resserrée sur trois paliers, force une progression visible avant que l'élan initial ne retombe.



Feuille de route 90 jours

0-30 jours : cartographier

Le premier palier ne consiste pas à corriger quoi que ce soit, il consiste à savoir. Établir un registre des usages IA existants, y compris ceux qui n'ont jamais été formellement validés (le « shadow AI » du Chapitre 3), et identifier les architectures RAG et les connecteurs de type MCP déjà en place. Dans la plupart des missions que je conduis, ce travail suffit à lui seul à révéler l'ampleur réelle de la dette algorithmique, presque toujours supérieure à ce que la direction générale imagine.

31-60 jours : sécuriser

Le deuxième palier cible les points les plus critiques identifiés au premier palier, pas l'ensemble du système d'information. Concrètement : cloisonner les accès des agents IA connectés à des systèmes sensibles (CRM, ERP, données personnelles), mettre en place une journalisation minimale là où elle n'existe pas, et documenter les sources utilisées par les systèmes RAG les plus exposés. L'objectif n'est pas l'exhaustivité, c'est la réduction rapide du risque le plus visible.

61-90 jours : structurer

Le troisième palier pose les bases organisationnelles développées aux Chapitres 8 et 9 : une première matrice de responsabilité (qui porte quelle dimension), et un premier niveau de maturité mesuré, même approximatif, pour chacune des cinq dimensions du cadre STAIR. Ce n'est pas un audit complet, c'est un point de référence à partir duquel progresser.

Ce que cette feuille de route n'est pas

Elle ne remplace pas une mise en conformité complète avec l'AI Act, le RGPD ou un référentiel comme ISO/IEC 42001 (Chapitre 8). Elle ne traite pas non plus, en 90 jours, une dette accumulée parfois depuis plusieurs années. Je l'ai conçue comme un point de départ délibérément resserré, destiné à produire une preuve tangible de progrès avant le prochain comité de direction, condition souvent nécessaire pour obtenir les moyens d'aller plus loin.

Conclusion du chapitre

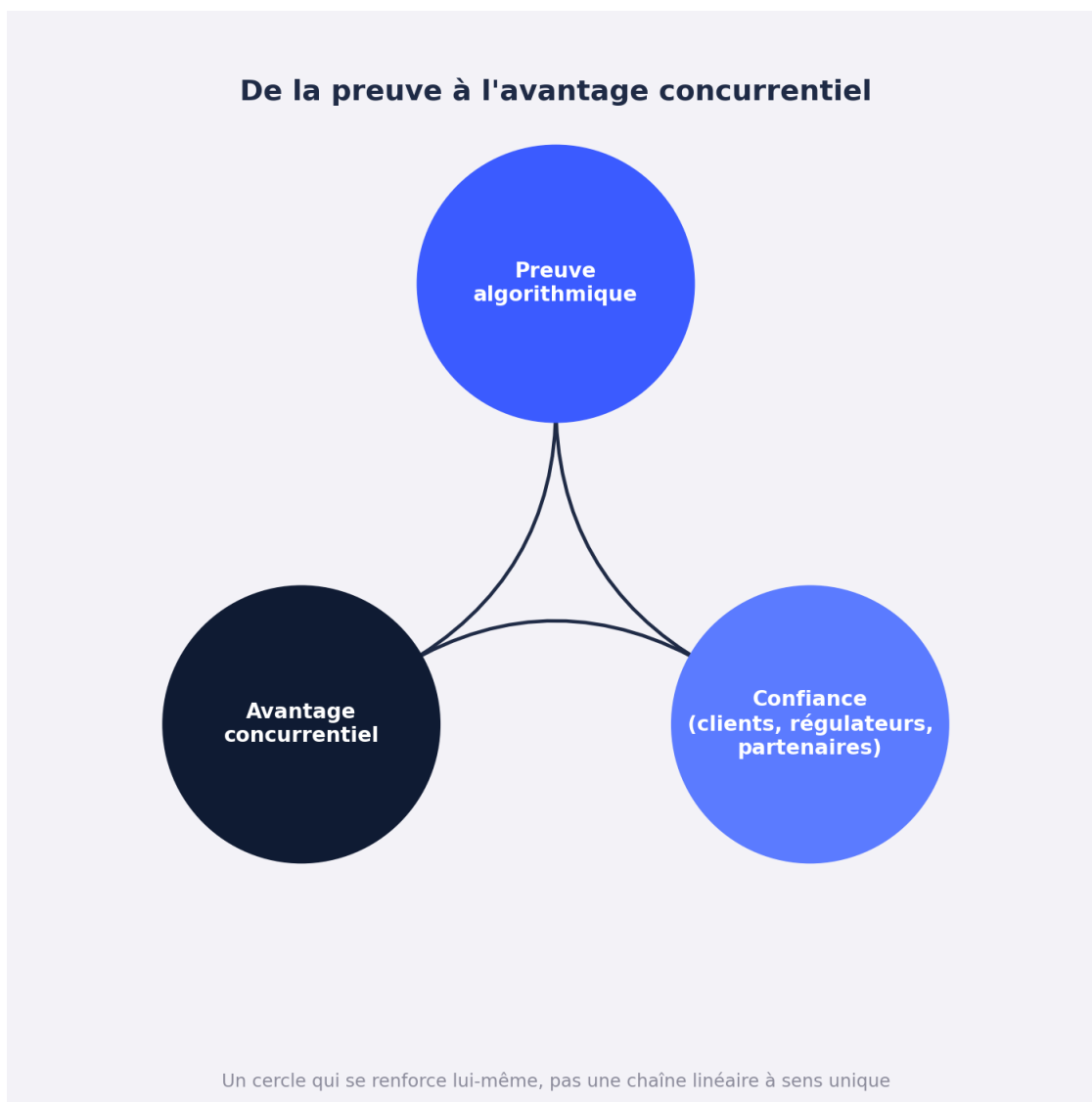
La feuille de route ne demande pas un budget de transformation. Elle demande une décision de démarrer, et trois paliers suffisamment courts pour que cette décision reste visible dans le temps. Le dernier chapitre de ce livre blanc referme la boucle ouverte au Chapitre 1 : pourquoi cette démarche, une fois engagée, cesse d'être une contrainte pour devenir un avantage.

Chapitre 11 : De la preuve à l'avantage concurrentiel

PARTIE 5 : ACTION

Ce livre blanc s'est ouvert sur un constat : l'IA générative a quitté le pilote, elle est diffuse dans le système d'information, et la plupart des organisations ne savent pas précisément ce qu'elle y fait. Il se referme sur une idée simple, que je défends dans chacun des dix chapitres précédents : la gouvernance de la preuve algorithmique n'est pas seulement une protection contre un risque, elle est un avantage que très peu d'organisations ont encore construit.

Un cercle qui se renforce lui-même



De la preuve à l'avantage concurrentiel

La preuve algorithmique, traçabilité, accès maîtrisé, sourcing vérifié, réversibilité, intégration technique gouvernée (Chapitre 8), produit de la confiance : celle des clients qui savent ce qu'une entreprise fait de leurs données, celle des régulateurs qui trouvent une documentation prête au moment d'un contrôle (Chapitre 5), celle des partenaires qui peuvent auditer une architecture RAG ou MCP sans découvrir un angle mort (Chapitre 4). Cette confiance, à son tour, devient un avantage concurrentiel mesurable : accès à des marchés ou des clients plus exigeants, cycles de vente raccourcis, réduction du coût du capital réglementaire. Et cet

avantage justifie de nouveaux investissements dans la preuve, ce qui referme et renforce le cercle.

C'est l'inverse du raisonnement encore dominant dans la plupart des comités de direction, où la gouvernance est perçue comme un coût qui s'ajoute à la valeur produite par l'IA. Les données que j'ai réunies dans ce livre blanc, l'écart de valeur documenté par le BCG (Chapitre 2), le coût chiffré du shadow AI par IBM (Chapitre 3), les litiges déjà tranchés en 2025 (Chapitres 2 et 3), montrent que c'est l'absence de preuve qui coûte, pas sa construction.

Ce que ce livre blanc a démontré, et ce qu'il n'a pas démontré

Par honnêteté envers un lectorat de dirigeants habitués à trancher entre signal et bruit : je n'ai pas prouvé que toutes les organisations gouvernées surperforment déjà celles qui ne le sont pas, les données disponibles ne permettent pas cette affirmation à ce stade (Chapitre 2). J'ai montré, avec des sources vérifiables et datées, que l'absence de preuve a déjà un coût mesurable dans les cas où elle se matérialise, que ce coût progresse avec l'intensification du cadre réglementaire (Chapitre 5), et que les organisations les plus avancées dans la capture de valeur IA sont aussi, structurellement, celles qui ont le plus investi dans la maîtrise de leurs systèmes (BCG, Chapitre 2).

Le point de départ, une dernière fois

J'ai montré au Chapitre 6, à partir d'une situation que j'ai vécue de l'intérieur, que ce problème n'est pas nouveau : hier des fichiers Excel non gouvernés, aujourd'hui des systèmes IA non gouvernés, la même cause organisationnelle, le manque de moyens ou de volonté de changement (Chapitre 3), produit la même dette. Les organisations qui ont déjà traversé cette transition une fois, avec leurs données et leurs systèmes d'information, ont un avantage qu'elles ignorent parfois : elles savent déjà à quoi ressemble le chemin. Il ne reste qu'à le refaire, plus vite, pour l'IA.

À propos de l'auteur

Je m'appelle Réналd Féré. Depuis quinze ans, j'accompagne des dirigeants sur la stratégie digitale, les systèmes d'information, la data et l'intelligence artificielle.

Mon métier n'a jamais eu de titre standard, et c'est précisément ce qui en fait l'intérêt. Je ne suis ni un DSI pur ni un directeur marketing pur : je suis le pont entre les deux. C'est une position qui permet de voir des problèmes que les silos ne voient pas, et c'est de là que vient tout ce livre blanc.

Concrètement, mon terrain couvre la stratégie digitale et le positionnement, l'architecture CRM et ERP, l'acquisition (SEO, SEA, et aujourd'hui la visibilité dans les moteurs génératifs), la gouvernance de la donnée, et l'intégration de l'IA dans les process métier, y compris ses briques techniques : architectures RAG, agents connectés via MCP, automatisations.

J'ai piloté plus de 50 millions d'euros de budgets digitaux et marketing, contribué à lever 7 millions d'euros pour des projets de diversification, et dirigé pendant six ans la stratégie digitale et marketing de la Fondation 30 Millions d'Amis, à l'échelle nationale.

Aujourd'hui, j'interviens en direction digitale à temps partagé, ce qu'on appelle le CDO part-time ou la direction fractionnelle. Le principe est simple : une organisation qui a besoin d'une direction digitale et SI expérimentée, mais pas d'un poste à temps plein, peut en disposer quelques jours par mois, avec la responsabilité et l'engagement d'un membre de l'équipe dirigeante.

Parlons de votre système d'information

Si ce livre blanc vous a donné envie de savoir où en est votre organisation, la meilleure façon de le découvrir n'est pas de remplir seul une grille d'auto-évaluation.

Je propose un premier échange de trente minutes, gratuit et sans engagement, pour comprendre votre contexte et identifier vos deux ou trois priorités réelles. Nous y regardons ce qui tourne déjà en IA chez vous, souvent plus de choses que ne le pense la direction générale, où se loge votre dette algorithmique, et par quel palier de la feuille de route à 90 jours il est le plus utile de commencer.

Et si votre sujet est plus large que l'IA, refonte d'un ERP, structuration d'un CRM, remise à plat d'une stratégie digitale, gouvernance de la donnée, pilotage de prestataires, c'est exactement mon terrain. N'hésitez pas à me solliciter, même si votre question ne rentre dans aucune case.

Réserver un échange stratégique de 30 minutes : www.reffeai.com/rendez-vous

Vous pouvez aussi me joindre sur LinkedIn, où je publie régulièrement sur ces sujets : [linkedin.com/in/renaldferé](https://www.linkedin.com/in/renaldferé)

Et si ce document vous a été utile, le meilleur service que vous puissiez me rendre est de le transmettre à la personne de votre organisation qui devrait le lire.