



CIO.FOCUS

**Comment les DSI encadrent
les risques liés à l'IA**

EN BREF

Souvent sous la pression des directions générales, les DSI déploient de plus en plus des IA en production. Or, ces technologies introduisent des risques nouveaux, notamment quand on parle de GenAI. Comment les DSI parviennent-ils à faire passer ce discours en interne ? Comment les DSI arbitrent-ils entre bénéfices et risques ? Quelles nouvelles structures et outils déploient-ils pour encadrer ces risques, pendant les phases de développement, mais aussi en production ? A quelles difficultés se heurtent-ils ?

Pour toute demande concernant CIO.focus :
contact-cio@it-news-info.com

Une publication de IT NEWS INFO :
14 Bd Poissonnière 75009 Paris

Rédacteur en chef :
Reynald Flechaux
reynald.flechaux@it-news-info.com

Principaux associés :
IT Facto et International Data
Group Inc.

Président et Directeur de publication :
Nicolas Beaumont

Directeur général : Nicolas Beaumont

CIO est édité par IT NEWS INFO,
SAS au capital de 3 000 000 €

Siret : 500034574 00029 RCS Nanterre

SOMMAIRE

/ STRATÉGIE

GenAI : les DSI tentent de garder le contrôle...
pour ne pas payer les pots cassés

3

/ TECHNOLOGIES

A Cannes, la gouvernance de l'IA reste
en bas des marches

7

/ CONTENU SPONSORISÉ

IA et fonctions finance : la vitesse seule
ne suffit pas, la fiabilité reste essentielle

10

/ STRATÉGIE

A la recherche d'une gouvernance
efficace de l'IA

13

/ STRATÉGIE

Les risques liés à l'IA au coeur
des préoccupations des entreprises

16

/ CONTENU SPONSORISÉ

Gestion des notes de frais : comment l'IA
renforce la fiabilité et la fluidité

18

/ TECHNOLOGIES

L'arrivée de l'IA dans le développement
d'applications critiques secoue les DSI

21

/ CONTENU SPONSORISÉ

Comptabilité fournisseurs : l'IA accélère,
le contrôle humain reste indispensable.

26

/ STRATÉGIE

Qui sera le premier DSI licencié à cause
des agents IA ?

29

/ TECHNOLOGIES

Évaluation et tests : le coût caché
du déploiement des agents IA

32



/ STRATÉGIE

GenAI : les DSI tentent de garder le contrôle... pour ne pas payer les pots cassés

L'engouement pour la GenAI ressemble à la ruée vers l'or. Métiers comme directions générales poussent à la roue. Dans ce contexte, les DSI se doivent de calmer les ardeurs pour éviter ou, a minima, contrôler les nombreux risques inhérents à cette technologie.

© istock



© DR

« Sur 200 cas d'usage proposés par les métiers, environ 20 ont été mis en production pour le moment », dit Lionel Chaine, DSI Bpifrance.

Aujourd'hui, il n'est pas possible d'éviter le sujet de l'IA générative », constate Lionel Chaine, DSI Bpifrance. A la Maif, « quand cette technologie a émergé, le sujet a été porté au niveau de la direction générale », renchérit Nicolas Siegler, DGA adjoint et DSI de l'assureur. Un engouement bien sûr dopé par les fournisseurs qui s'assurent ainsi des gains substantiels même si, sur le terrain, peu d'entreprises constatent un retour sur investissement quantifiable. Une étude menée par le MIT (The GenAI Divide: State of AI in Business 2025) estime que 95 % des entreprises ne constatent aucun retour mesurable de leurs investissements en IA générative. « Pour les cas d'usage pertinents, les bénéfices émergent seulement après la mise à l'échelle », souligne Lionel Chaine.

Mais même cette absence de visibilité n'empêche pas la quasi-totalité des directions dans les entreprises, administrations ou collectivités de tester massivement cette technologie. Une autre raison explique un tel emballement, à savoir la banalisation de l'utilisation de ChatGPT et autres LLM par les particuliers. « Même quand une entreprise interdit le Shadow AI, des salariés sont tentés de le contourner en passant par leur smartphone ou un abonnement personnel », observe Damien Peries, DSI de Crosscall, un fabricant de smartphones tout-terrain. Les responsables IT se doivent donc de suivre le mouvement en contrôlant au mieux les risques liés aux usages de cette technologie. Une démarche d'autant plus nécessaire qu'ils pourraient être considérés comme responsables des

déboires liés à l'IA générative. Et ceux-ci se déclinent sur de nombreux points, de la réglementation aux hallucinations en passant par l'exposition de données sensibles, les impacts RH, environnementaux ou encore les dérapages budgétaires. Le Cigref a publié récemment plusieurs guides sur le sujet, [dont un dédié à la mise en oeuvre pratique de l'AI Act européen](#).

Sur le terrain, les DSI jonglent donc pour prendre en compte et pondérer tous ces facteurs en fonction de leur organisation. Une constante, cette démarche passe ou devrait passer par la mise en place d'une vraie gouvernance et ce, à l'échelle de l'organisation. « C'est la bonne approche pour bénéficier des apports de cette technologie en maîtrisant au mieux les risques induits, confirme Nicolas Siegler, DGA et DSI de la Maif. Dans ce but, nous avons créé un conseil de surveillance 'Numérique éthique' interne pour évaluer les impacts RH, environnementaux et éthiques, liés à ces nouveaux usages ».



Pour bénéficier des apports de cette technologie en maîtrisant au mieux les risques induits, nous avons créé un conseil de surveillance 'Numérique éthique' interne »

Nicolas Siegler

A la recherche du ROI de l'IA

La gouvernance commence par une étape qui n'a rien de spécifique à l'IA, à savoir une évaluation de l'apport de cette technologie. Celle-ci reste complexe parce qu'elle dépend du type de cas d'usage, intégré dans un processus métier ou plus générique. Dans le premier cas de figure, le retour sur investissement reste relativement simple à évaluer, même si la décision finale de mise en production dépend également d'autres facteurs. Chez Bpifrance, « sur 200 cas d'usage proposés par les métiers, environ 20 ont été mis en production pour le moment », illustre Lionel Chainé. Pour le deuxième cas, qualifié d'IA horizontale dans le rapport du Cigref, qui couvre des tâches comme la génération d'un compte rendu de réunion, le résumé d'un dossier, la réponse à des questions techniques,



© Mélanie Chaigneau/MAIF

« Pour 'surveiller' le chatbot chargé de répondre à des questions sur la prise en charge des incidents dans les contrats, une équipe dédiée de deux postes a été créée », note Nicolas Stiegler, DGA adjoint et DSI de Maif.

plusieurs études citées par le rapport estiment à 50 minutes le gain de temps quotidien par utilisateur. Mais, en fonction des organisations, le retour sur investissement n'est pas seulement financier. A Lorient, un chatbot allège les tâches des agents d'accueil de l'agglomération sans qu'il soit question de gagner du temps, mais plutôt de soulager la charge de travail de ces agents. Problème, dans ses versions ouvertes à tous (ChatGPT...), cette IA horizontale est aussi utilisée par une part des salariés dans un contexte très métier. Des pratiques qui font courir des risques à leurs organisations et imposent, pour les DSI, de supprimer ou à minima de contrôler cette Shadow IA.

Sensibiliser aux risques plutôt qu'interdire

Beaucoup d'organisations ont opté pour une sensibilisation de leurs collaborateurs. Démarche d'autant plus nécessaire que prendre des solutions payantes, comme Copilot pour l'ensemble des salariés est irréaliste, soulignent Crosscall et la Maif, entre autres. Chez Crosscall, « tout ce qui n'est pas listé aux collaborateurs comme autorisé est interdit, insiste Damien Peries. Le but n'est pas d'empêcher toute requête, les flux ne sont pas bloqués, mais de responsabiliser les utilisateurs par des actions différentes de sensibilisation. » D'autres organisations vont plus loin dans l'accompagnement de leurs salariés.



La Maif autorise l'utilisation de Copilot web sur Azure, une version gratuite de Copilot, pour tous ses collaborateurs, « ce qui sécurise ces usages », précise Nicolas Siegler. Si l'accès à d'autres LLM n'est pas techniquement bloqué, tous les flux sont monitorés et les utilisateurs sensibilisés aux risques associés. L'assureur a mis en place une cellule d'assistance aux prompts. « Bpifrance a également mis en place une charte d'utilisation », décrit Lionel Chaine. DSI de l'agglomération de Lorient, Alain Cottencin a choisi la même approche non coercitive : « nous privilégions

la sensibilisation et la pédagogie. Une note a été diffusée à l'ensemble des agents pour les informer des opportunités et des risques liés à ces usages. Une démarche qui concerne également les directions ». Chez Crosscall, « le management l'utilise surtout pour générer des comptes-rendus de réunion ou chercher des informations dans ces documents, détaille Damien Peries. La sensibilisation leur a permis de prendre conscience des risques à exposer des informations sensibles, une simple proposition commerciale, par exemple. » En d'autres mots, à ne pas utiliser ChatGPT dans un contexte professionnel. Si l'acculturation n'en est qu'à ses débuts, la Shadow IA la plus difficile à éviter vient de certains éditeurs « imposant l'IA générative dans les montées de versions », signale Nicolas Siegler.



« Nous privilégions la sensibilisation et la pédagogie. Une note a été diffusée à l'ensemble des agents pour les informer des opportunités et des risques », souligne Alain Cottencin, DSI de l'agglomération de Lorient.

Surveiller les hallucinations

Outre l'exposition des données sensibles, cet accompagnement a également pour but de faire prendre conscience que cette technologie est loin d'être parfaite et génère des hallucinations. « Il est indispensable que l'utilisateur garde un regard critique », résume un DSI du secteur de la santé. Et se reposer sur ce seul filtre n'est pas suffisant. Le Cigref préconise de mettre en place une organisation ad hoc. Concrètement, il s'agit de définir une matrice

définissant les rôles et responsabilités en interne sur les cas d'usage critiques.

Chez Bpifrance, le recours aux LLM repose sur des agents IA classés en trois catégories. Les premiers crawlent le web, mais ne peuvent accéder aux données internes. Le second niveau autorise un accès hybride, mais limité aux données internes non critiques. Tandis que le troisième se focalise sur les seules données internes. « Pour limiter les risques d'exposition de données et plus globalement les risques cyber, nous nous interdisons d'utiliser de la GenAI pour les cas d'usage critiques. Comme par exemple, celui chargé de l'octroi de prêts, qui fonctionne actuellement à partir de règles et de machine learning. Il est possible que nous utilisions la GenAI pour ces cas à l'avenir si des solutions robustes permettent de pallier les risques », décrit Lionel Chaine.

« Pour limiter les risques d'exposition de données et plus globalement les risques cyber, nous nous interdisons d'utiliser de la GenAI pour les cas d'usage critiques. »

Lionel Chaine

Côté Maif, un dispositif baptisé Maintien en Conditions Intelligentes a été mis en place pour vérifier les éventuels dérapages. Il prend différentes formes selon les cas d'usage. « Pour 'surveiller' le chatbot chargé de répondre à des questions sur la prise en charge des incidents dans les contrats, une équipe dédiée de deux postes a été créée. Pour un autre chatbot, encore en expérimentation interne et qui sera chargée de répondre aux réclamations de nos sociétaires, les collaborateurs testeurs contrôlent régulièrement les éventuels dérapages », décrit Nicolas Siegler.

Pas de firewall pour les IA génératives

Le risque ne se limite pas au mécontentement de certains clients ou à une dégradation de l'image de l'entreprise. Il se décline également sur le volet réglementaire, avec des sanctions financières conséquentes. L'AI Act

impose aux entreprises une transparence sur l'usage de cette technologie ou encore un suivi des données. Chez Bpifrance, « tous les usages sont supervisés, les prompts sont archivés associés au profil de l'utilisateur », décrit Lionel Chaine. La Maif a développé une plateforme qui gère tous les flux liés à l'utilisation de LLM, « ce qui permet de suivre les impacts environnementaux », signale en passant Nicolas Siegler. Autre risque, en tant qu'entité publique, la DSI de l'agglomération de Lorient se doit d'éviter une exposition d'une part de ses données. Si la DSI teste la GenAI avec les solutions phares du marché, elle envisage de mettre cette technologie en production avec des modèles Open Source ou un fournisseur européen, comme Mistral.

Côté cyber, éviter d'exposer les données sensibles reste nécessaire. « À ce jour, les équivalents d'un firewall pour les IA génératives ne sont pas encore au point, par exemple pour éviter l'injection de prompts », rappelle Lionel Chaine. Même l'utilisation d'outils sur étagère censés protéger les données ne met pas les organisations totalement à l'abri. « Nous surveillons le dark web pour vérifier que des données de nos sociétaires ne sont pas divulguées », souligne ainsi Nicolas Siegler. Au final, comme le dit le rapport du Cigref, « l'appréhension des risques doit être faite globalement et non sur chaque facteur ». Sur le terrain, la confidentialité, la sécurité, la protection des données, les hallucinations et la propriété intellectuelle sont souvent imbriquées. Un numéro d'équilibriste pour les DSI, qui n'ont d'autre choix que de suivre le mouvement.



UN ARTICLE RÉDIGÉ PAR

Patrick Brébion, journaliste

Après quelques années dans le développement logiciel, Patrick est devenu journaliste. Depuis le siècle dernier, il couvre plus spécialement tous les sujets IT pour la presse spécialisée. Il a également passé quelques années en tant qu'ingénieur de recherche à l'Université.

/ TECHNOLOGIES

A Cannes, la gouvernance de l'IA reste en bas des marches

Les débats lors du World AI Cannes Festival, qui se tenaient en février 2026, soulignent le chemin qui reste à parcourir en matière de régulation des usages de l'IA en entreprise. La conformité à la législation étant loin de résumer les défis qui attendent DSI et Chief Data Officer.



De gauche à droite, Nadia Filali, directrice innovation du groupe Caisse des dépôts, Laura Grassi, de l'Ecole polytechnique de Milan, et Prad Sharma, directeur des technologies émergentes de Citigroup, lors d'une table ronde au World AI Cannes Festival.

Inventer la gouvernance de l'IA. C'est un des défis qui attend les entreprises, qui voient la technologie évoluer bien plus rapidement que leur capacité à en encadrer les usages, comme l'ont montré les débats lors du World AI Cannes Festival, qui se déroulait dans la ville balnéaire, les 12 et 13 février. Le tout dans un environnement réglementaire encore incertain et, surtout, fragmenté. « Les Etats-Unis, la Chine et l'Europe ont des philosophies très différentes en matière de régulation », explique Oliver Patel, responsable de la gouvernance de l'IA au sein du groupe pharmaceutique AstraZeneca. Ce dernier souligne encore que « rien d'aussi exhaustif et structuré que l'AI Act n'existe en dehors de l'Europe. » Et ce, même si les ambitions initiales du texte ont été récemment revues à la baisse.

« La France porte une vision unique d'une IA responsable, durable, au service du bien commun qui va devenir un atout compétitif », défend Anne Le Hénanff, ministre de l'IA et du numérique, qui s'exprimait en ouverture du salon cannois, à quelques jours du sommet international sur l'IA qui se tient cette semaine à New Delhi (Inde). Des visions politiques avec lesquelles se débattent les entreprises, en tout cas les multinationales. « Elles font face à des corpus réglementaires divergents », reprend Oliver Patel, qui souligne aussi que l'AI Act n'a pas créé d'effet domino dans d'autres parties du monde.

« Nous ne pouvons pas attendre la régulation »

Difficile dès lors de construire un cadre de gouvernance figé et unifié, d'autant que les frameworks publiés par des organismes comme l'OCDE ou l'Unesco ne descendent pas au niveau des règles opérationnelles. « Les multinationales ont besoin de politiques globales, associées à des mécanismes apportant de la flexibilité, pour tenir compte des besoins spécifiques à une entité. Nous devons décider des risques à prioriser tout en restant pragmatiques. Parfois le risque pris en n'adoptant pas l'IA est supérieur à celui qu'il nous fait courir », souligne Oliver Patel.

« *Nous devons décider des risques à prioriser tout en restant pragmatiques. Parfois le risque pris en n'adoptant pas l'IA est supérieur à celui qu'il nous fait courir.* »
Oliver Patel

« Nous ne pouvons pas attendre une coordination et une régulation globales », abonde Lisa Bechtold, directrice de la gestion des risques de Nestlé, pour qui ces hésitations des régulateurs peuvent aussi se muer en atout, une organisation ayant ainsi une certaine marge de manoeuvre pour affirmer ses positions vis-à-vis du marché. « Si un million de labels de confiance existent, qui s'en souciera ? lance la responsable du groupe suisse. Mieux vaut dès lors se concentrer sur une standardisation globale à l'échelle de l'organisation, plutôt que sur les régulations elles-mêmes. La gouvernance de l'IA ne doit pas être vue comme un fardeau administratif, mais comme un investissement dans la sécurité. »

Gouverner des milliers d'agents ?

Se saisir du sujet sans attendre est d'autant plus essentiel que l'arrivée des agents se traduit par une forme de délégation de l'autorité humaine à la machine sur certaines décisions. Posant de nouveaux défis en matière de gouvernance. « Tout part de la définition des risques associés à chaque agent, qui

va permettre de préciser le rôle de l'humain dans sa supervision », dit Lisa Bechtold. Mais ce principe clef de la gouvernance se heurte à une question de volumes.

« *Tout part de la définition des risques associés à chaque agent, qui va permettre de préciser le rôle de l'humain dans sa supervision.* »
Lisa Bechtold

« Certes, conserver un contrôle humain permet de prévenir les dérives des modèles et de mettre en place des validations dans des processus automatisés, dit Oliver Patel. Mais ce principe dit 'human in the loop' est bousculé avec l'IA agentique, car il ne s'applique plus pour des systèmes susceptibles de prendre des milliers de décisions par seconde. » Même préoccupation pour Lisa Bechtold, qui en plus d'une approche basée sur les risques, recommande d'investir dans la formation des collaborateurs et de mener un travail technique sur la cohérence des protocoles spécifiques aux agents. Car, pour l'heure, les praticiens en entreprises ne peuvent guère s'appuyer sur des normes ou standards internationaux, encore inexistantes à l'exception d'un cadre de gouvernance tout juste publié par Singapour.

Encore une fois, les déploiements paraissent donc devancer les pratiques de gouvernance. « En 2026, si vous en êtes encore aux expérimentations, vous êtes en retard, assure Prad Sharma, le directeur des technologies émergentes de Citigroup. L'IA agentique va fonctionner 24 heures sur 24 au sein de vos systèmes ; vous devez trouver un moyen de la monitorer. » Pour cet expert, la gouvernance de cette technologie sera bien le défi de 2026. « Pour intégrer l'IA dans des usages à haut risque, vous devriez comprendre de quoi les modèles sont faits, donc comment ils sont entraînés, souligne-t-il. Et quand des modèles sont intégrés par les éditeurs dans vos outils standards, vous faites face à un autre risque. Et bientôt ce sera également le cas au coeur même des systèmes d'exploitation. Chacun de ces modèles est un point de défaillance potentiel. » Et de souligner que cette intégration des modèles d'IA - LLM en tête - par les éditeurs va soulever des questions de responsabilité, en cas de résultat inapproprié.

Tous les LLM sont biaisés

Tous ces algorithmes arrivant en entreprise pré entraînés comportent en effet, par définition, des biais. « Avec l'IA générative, vous n'êtes jamais sûr que les données d'entraînement soient représentatives, insiste Laura Grassi, professeure associée et responsable de l'observatoire Fintech & Insurtech de l'Ecole polytechnique de Milan. Le risque est d'autant plus grand que de nombreuses organisations utilisent les mêmes modèles concentrant les mêmes biais (une large majorité des données d'entraînement émanant de quelques pays). Avec des conséquences très concrètes, comme le souligne par exemple Ieshika Chandra, une Data Scientist passée par EY et Lyft et actuellement chez Walmart. « Sur les sites de commerce électronique des distributeurs, ce biais peut entraîner une surreprésentation des acteurs dominants (comme Nike sur le marché des baskets). Les retailers doivent intégrer davantage de fonctionnalités orientées vers l'explication des résultats afin de renforcer la confiance des consommateurs et des vendeurs », explique Ieshika Chandra.

« Si les données les moins sensibles peuvent être confiées à des LLM tiers, nous avons besoin de davantage de contrôle sur les informations en accès restreint. »

Prad Sharma

Face à des modèles par nature opaques, la meilleure garantie pour les entreprises réside dans le contrôle de la data, la gouvernance de l'IA prolongeant celle de la data. « Si les données les moins sensibles peuvent être confiées à des LLM tiers, nous avons besoin de davantage de contrôle sur les informations en accès restreint, dit Prad Sharma, de Citigroup. Si vous n'avez pas classé vos données selon différents niveaux de sensibilité, vous n'êtes tout simplement pas prêts ! ». Une nécessité qui passe aussi par une sensibilisation des métiers à ces questions. C'est le pari que fait la Caisse des dépôts avec son programme Docuscore, inspiré du Nutriscore dans l'alimentaire. « L'objectif est de motiver les collaborateurs à mieux gérer les documents et l'information », résume Nadia Filali, directrice innovation du groupe public.

Des biais, mais vos biais

Les LLM sont par nature biaisés. L'enjeu n'est donc non pas d'en supprimer les biais, mais d'en ajouter. C'est en somme le message – « contrintuitif », comme il le reconnaît lui-même – de Luc Julia, ex-directeur de la stratégie scientifique de Renault parti monter une start-up (LucStunts) à l'ambition encore floue, lors du World AI Cannes Festival. « Les LLM sont majoritairement américains et répondent comme un bon étudiant américain », explique l'ingénieur, qui cite notamment le cas de l'invention de l'avion que les chatbots attribuent souvent aux frères Whright et non à Clément Ader.



Luc Julia, lors du World AI Cannes Festival, le 13 février.

« Espérer enlever les biais de milliards de données est illusoire. Il faut assumer au contraire d'introduire d'autres biais en injectant des données françaises, africaines, singapouriennes ou que sais-je ! » Selon lui, c'est même le contrôle de ces biais qui donne à une IA son efficacité. « Un agent, c'est justement une IA spécialisée faisant les choses correctement parce qu'elle est biaisée sur un domaine spécifique, avec des données dans lesquelles vous avez confiance », souligne Luc Julia.



UN ARTICLE RÉDIGÉ PAR

Reynald Fléchaux, Rédacteur en chef de CIO

Suivez l'auteur sur Linked In



/ CONTENU SPONSORISÉ PAR MEDIUS

IA et fonctions finance : la vitesse seule ne suffit pas, la fiabilité reste essentielle

Derrière les gains de productivité apportés par l'IA aux fonctions finance, de nouveaux enjeux émergent. Pour les DAF et les DSI, la question n'est plus seulement de déployer l'IA, mais de structurer son utilisation pour qu'elle soit fiable, explicable et maîtrisée. Ahmed Fessi, Directeur de la Transformation et des Systèmes d'Information de Medius, livre son analyse et des pistes d'action pour ancrer durablement l'IA dans les processus financiers.

© Istock



© DR

Ahmed Fessi, Directeur de la Transformation et des Systèmes d'Information de Medius.

L'IA intervient aujourd'hui tout au long du cycle finance, y compris dans le traitement des factures fournisseurs et des notes de frais. Elle constitue un puissant outil d'aide à la décision, à condition d'être mobilisée avec discernement. L'IA peut générer des résultats très plausibles, mais pas toujours corrects. « C'est précisément pour cela que l'expertise humaine reste essentielle », explique Ahmed Fessi.

Pour ce dernier, l'IA renforce le rôle des collaborateurs plutôt qu'elle ne le diminue : « l'humain est la première ligne de défense, mais aussi la première source de valeur. Il conserve pleinement sa responsabilité et sa capacité d'arbitrage. L'IA lui apporte des éléments supplémentaires pour décider plus vite et plus efficacement, sans jamais se substituer à son jugement. » Cette vigilance est particulièrement importante lorsque l'IA intervient dans des décisions sensibles, comme la détection de fraudes internes. L'IA génère un éclairage précieux, mais le raisonnement humain demeure indispensable pour valider et contextualiser les résultats.

Vers une IA explicable et gouvernée

Pour ancrer durablement l'IA dans les processus financiers, la compréhension des modèles devient un facteur clé de confiance.

« [L'IA repose sur des modèles probabilistes](#), là où l'automatisation classique suit des règles strictes et déterministes. Concrètement, l'IA analyse les données et identifie des schémas pour proposer des

recommandations ou des prédictions, plutôt que de produire systématiquement le même résultat », détaille Ahmed Fessi.

Ahmed Fessi estime que « l'enjeu n'est pas de décortiquer chaque ligne de code, mais de garantir que le modèle est robuste, traçable et utilisé dans des conditions fiables, afin que les décisions, sur lesquelles il apporte un éclairage, restent cohérentes et maîtrisées. »

Cette exigence d'explicabilité constitue un levier de maturité pour les organisations. Elle permet non seulement de sécuriser les décisions, mais aussi de mieux les documenter et de renforcer leur légitimité interne et externe.

Chez [Medius](#), cette exigence se traduit par la publication de « model cards » qui documentent les données d'entraînement, les usages prévus, les limites des modèles et les tests réalisés. « L'objectif est de donner à nos clients une vision claire du fonctionnement général du modèle et des garanties associées », précise Ahmed Fessi.

Au-delà de la performance technologique, cette transparence s'inscrit également dans un cadre réglementaire renforcé avec le EU AI Act. Plus qu'une contrainte, ce cadre contribue à professionnaliser l'usage de l'IA et à structurer sa gouvernance au sein des directions financières.

Une bonne décision repose sur des données fiables

Une décision explicable n'est fiable que si elle repose sur des données elles-mêmes fiables. Des informations bien structurées et pertinentes permettent à l'IA de produire des résultats précis et exploitables.

« Assurer la qualité des données est un travail rigoureux, mais essentiel », souligne Ahmed Fessi. « Même 1 % de données incorrectes peut être amplifié par l'IA, générant jusqu'à 5 ou 10 % d'erreurs dans les résultats », ajoute-t-il. C'est pourquoi il est important de nettoyer et de valider les données dès le départ. »

La qualité des données d'entraînement est aussi un levier fondamental pour la fiabilité de l'IA. « Plus les données utilisées pour entraîner le modèle sont précises et complètes, plus les résultats générés seront fiables », complète-t-il. « Si les données initiales ne sont pas

suffisamment rigoureuses, le modèle reproduira des erreurs. C'est pour cela que le contrôle sur la qualité des données en amont de la chaîne reste fondamental.

Et l'IA s'améliore grâce à un entraînement continu. « En intégrant régulièrement le feedback des utilisateurs, nous renforçons progressivement la fiabilité du modèle et la pertinence de ses recommandations », conclut le Directeur de la Transformation et des Systèmes d'Information de Medius.

Sensibiliser pour fiabiliser l'usage de l'IA

Même avec des données de qualité, la fiabilité peut être compromise si les outils eux-mêmes ne sont pas sécurisés. Selon PwC, [la cybercriminalité représente la troisième préoccupation pour les DAF](#) en 2026. L'utilisation de l'IA offre de nouvelles surfaces d'attaque. Les entreprises sont exposées à des menaces, comme les fuites de données ou l'usage non encadré de certains outils. Ce phénomène, parfois appelé Shadow AI, se produit lorsque les collaborateurs utilisent des solutions IA non validées ni supervisées par la DSI.



Pour Ahmed Fessi, le remède au Shadow AI n'est pas d'interdire l'utilisation de l'IA mais de former les salariés. « Le Shadow AI est la résultante naturelle d'une interdiction qui ne propose pas d'alternative. La vraie réponse, c'est donner les bons outils IA aux collaborateurs », estime-t-il. La formation est d'ailleurs très demandée : **9 collaborateurs sur 10 souhaitent être formés** selon une étude de KPMG.

La formation est cruciale car la DSI ne peut pas contrôler tous les faits et gestes des salariés. C'est ce que rappelle Ahmed Fessi : « il est tout à fait possible qu'un collaborateur utilise son téléphone personnel pour poser une question à un outil d'IA avec des données confidentielles. Ce type de situation échappe aux systèmes d'information. C'est pourquoi la sensibilisation reste, selon moi, la première mesure à mettre en place pour garantir un usage sûr et maîtrisé », conclut-il.

Quand la fraude met la confiance à l'épreuve

La fiabilité des décisions financières repose non seulement sur la sécurité des outils, mais aussi sur la capacité à détecter et prévenir les fraudes, y compris celles générées par l'IA. En effet, avec l'IA générative et l'évolution des modèles, certaines fraudes se complexifient. Avec l'IA générative, il est possible de créer rapidement de faux e-mails ou fausses factures, et les deepfakes rendent l'arnaque au président plus difficile à détecter en reproduisant la voix ou l'image de manière très réaliste. « Dans une étude menée auprès de nos clients, [53 % des directions financières](#) ont été confrontées à une fraude via un deepfake », signale Ahmed Fessi.

Pour lutter contre ces menaces, l'IA peut elle-même devenir un puissant outil de détection. « Elle permet d'identifier des comportements récurrents et d'analyser de larges volumes de données, en les comparant à des historiques pour repérer les anomalies potentielles », explique Ahmed Fessi.

Pour détecter la fraude commise par l'IA, la solution peut être l'IA elle-même. Ahmed Fessi en décrit le fonctionnement : « l'IA permet de détecter certains patterns et d'analyser de gros volumes de données puis de les comparer à des données historiques ».

Un usage de l'IA qui séduit de nombreuses organisations. « Les directions financières adoptent nos solutions non seulement pour automatiser leurs processus, mais aussi pour détecter les risques de fraude, que ce soit dans la comptabilité fournisseurs ou pour les notes de frais », explique Ahmed Fessi. L'IA ne crée de valeur pour les fonctions finance que si les décisions qu'elle alimente restent fiables, explicables et maîtrisées.

L'IA, levier d'efficacité lorsqu'elle produit un réel impact

Ahmed Fessi rappelle un chiffre frappant, issu du Massachusetts Institute of Technology (MIT) : 95 % des projets pilotes en IA n'aboutissent pas. « Pour créer de la valeur, l'IA doit aussi être utile », insiste-il. Lorsqu'elle est correctement utilisée, elle accélère l'extraction de données, assiste les collaborateurs dans la prise de décision et renforce la fiabilité des processus financiers.

Chez Medius, l'intelligence artificielle intégrée à ses solutions d'automatisation de la comptabilité fournisseurs et de gestion des notes de frais repose sur des cas d'usage concrets, identifiés en amont à partir de demandes récurrentes et d'expérimentations. « Notre objectif est d'avoir un impact réel et immédiat pour les DAF », explique Ahmed Fessi.

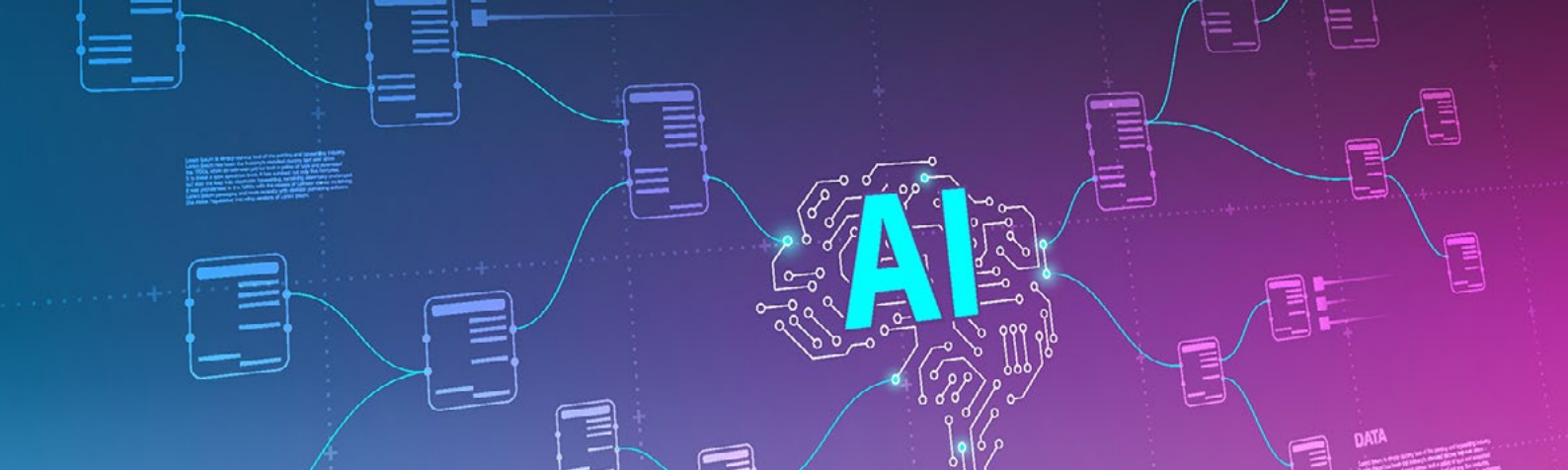
Cela peut se traduire par la fluidification du processus comptable, l'amélioration de la détection des fraudes et des erreurs, la prise en charge de tâches chronophages ou l'information des collaborateurs et fournisseurs, afin de recentrer les équipes sur l'analyse, le pilotage et la décision, et tendre vers une finance plus fluide et autonome.



UN ARTICLE RÉDIGÉ PAR

Medius

Contenu sponsorisé



/ STRATÉGIE

A la recherche d'une gouvernance efficace de l'IA

Selon les experts, la rapidité d'adoption de l'IA et l'évolution permanente des technologies complexifient la mise en oeuvre de pratiques de gouvernance au sein des organisations.

© istock

De nombreuses organisations déployant l'IA reconnaissent la nécessité de mettre en place des garde-fous, mais rares sont celles qui ont su élaborer un modèle de gouvernance mature. D'après une récente étude de Cisco, trois organisations sur quatre déclarent disposer d'un processus de gouvernance dédiés pour l'IA, mais seulement 12 % estiment que leurs efforts en la matière ont atteint un bon niveau de maturité.

L'étude de Cisco, *Data and Privacy Benchmark Study*, suggère non seulement que les processus de gouvernance de l'IA sont encore en plein développement, mais aussi que les préoccupations liées à la protection de la vie privée incitent les entreprises à renforcer les garde-fous. Ainsi, 93 % des organisations prévoient d'investir davantage pour faire face à la complexité croissante des systèmes d'IA et aux attentes des clients et des autorités de régulation.



Si tout le monde disait : "Oui, nous avons un programme, donc nous sommes prêts", ce serait faire preuve d'une certaine naïveté face à l'incroyable complexité du sujet. »

Jen Yokoyama

Pour Jen Yokoyama, vice-présidente senior en charge de l'innovation et de la stratégie juridique chez Cisco, voir les professionnels de l'informatique et de la sécurité interrogés reconnaître la nécessité de poursuivre leurs efforts en matière de gouvernance de l'IA est un signe encourageant. « Cette statistique illustre bien la prise de conscience de la complexité à laquelle ces entreprises sont confrontées, explique-t-elle. Si tout le monde disait : "Oui, nous avons un programme, donc nous sommes prêts", ce serait faire preuve d'une certaine naïveté face à l'incroyable complexité du sujet. »

L'adoption précède la gouvernance

L'un des principaux défis pour les organisations qui déploient l'IA réside dans le décalage entre adoption et gouvernance, souligne Jen Yokoyama. De nombreux responsables informatiques doivent prendre des décisions concernant la conformité, les questions d'éthique et la transparence au fur et à mesure du déploiement de la technologie, ajoute-t-elle. « Une grande partie du problème réside dans la volonté d'aller vite, d'adopter rapidement l'IA et de rentabiliser cette technologie, explique-t-elle. Il faut agir vite, car les gens investissent maintenant et veulent voir des résultats. Il faut donc trouver des solutions à un rythme soutenu. »

Les responsables IT doivent également prendre en compte des questions telles que la confidentialité, le partage des données avec les fournisseurs d'IA, ainsi que la localisation et la souveraineté des données lors du lancement de leurs projets, précise Jen Yokoyama. « Le problème, c'est qu'une solution unique ne convient pas dans tous les cas », ajoute-t-elle.

Les multiples facettes de la gouvernance

La rapidité d'adoption de l'IA a complexifié les efforts de gouvernance, confirme Jean-Matthieu Schertzer, directeur de l'IA chez Eagle Eye Group, fournisseur de solutions informatiques pour le marketing. « De nombreuses organisations déploient rapidement l'IA dans des fonctions telles que le marketing, pour des questions d'automatisation, de personnalisation et d'efficacité opérationnelle, mais la maturité de la gouvernance est souvent à la traîne à mesure que l'adoption se généralise, explique-t-il. L'opacité de nombreux systèmes d'IA rend difficile le suivi des décisions, l'identification des biais et l'établissement de responsabilités claires en cas de problème. »

Selon Jean-Matthieu Schertzer, une gouvernance efficace de l'IA repose sur des pratiques opérationnelles structurées, telles que la documentation des limitations des modèles, la réalisation d'audits de biais et de sécurité, et la mise en place de processus de revue et de supervision. Il ajoute que, parallèlement, les responsables de l'IA doivent répondre aux attentes croissantes en matière de transparence, de consentement et de conformité réglementaire, des enjeux qui

concernent de nombreux services au sein d'une organisation. « Ces exigences impliquent les équipes juridiques, data, sécurité, marketing et produit, et les progrès sont souvent ralentis lorsque les responsabilités ne sont pas clairement définies ou que les initiatives restent cantonnées à des projets pilotes cloisonnés au lieu de devenir des pratiques opérationnelles courantes », explique-t-il.

Une discipline émergente

Ron Davis, vice-président senior de l'ingénierie produit et responsable de l'IA chez QAD Redzone, éditeur de logiciels pour le manufacturing, ajoute que la rapidité des déploiements complexifie la gouvernance de l'IA, tout comme la vitesse d'évolution de la technologie elle-même. « L'innovation en IA progresse plus vite que la formalisation des contrôles dans la plupart des entreprises, obligeant les équipes à adapter simultanément la technologie et la gouvernance, déclare-t-il. Pour ne rien arranger, la gouvernance de l'IA demeure une discipline émergente, dont les normes et modèles opérationnels sont encore en cours d'élaboration. »

Le défi de la gouvernance de l'IA est amplifié dans le secteur manufacturier, car l'IA influence de plus en plus les décisions opérationnelles qui affectent la sécurité et la qualité, ajoute-t-il. « De ce fait, les organisations doivent composer avec l'incertitude tout en jonglant avec la rapidité, les risques, la responsabilité et la sécurité », explique Ron Davis.

Mauvaise gouvernance des données = IA mal gouvernée

Dans de nombreuses organisations, la difficulté à renforcer la gouvernance de l'IA découle d'un manque de gouvernance de la donnée elle-même, ajoute Anisha Vaswani, directrice des systèmes d'information et de la relation client chez Extreme Networks, fournisseur d'équipements réseau. De nombreuses entreprises peinent encore à mettre en place une gouvernance des données performante, constate-t-elle.

« À cela s'ajoutent l'évolution rapide des technologies, les nouveaux investissements et la nécessité de former le personnel pour une gouvernance efficace, ce qui ne fait qu'aggraver le problème, ajoute-t-elle. On

est confronté à une grande complexité et à une fragmentation des données et des modèles, ainsi qu'à une évolution rapide du paysage technologique de l'IA. Pour pouvoir la gouverner, il faut la comprendre, se tenir informé des évolutions. Or, celles-ci sont extrêmement rapides. »

Travail collaboratif

Anisha Vaswani recommande aux organisations de mettre en place des équipes transverses pour traiter les questions de gouvernance. Et les DSI et autres responsables IT ont un rôle crucial à y jouer pour promouvoir l'auditabilité et l'explicabilité des outils d'IA, ajoute-t-elle. « La gouvernance, c'est aussi adopter une attitude pessimiste et se demander : "Quels sont les risques et comment les atténuer ?", explique-t-elle. Nous n'avons pas encore trouvé de solution miracle ; il s'agit d'une technologie émergente. »

Jen Yokoyama, de Cisco, confirme que l'élaboration de bonnes pratiques nécessitera un travail collaboratif, car la gouvernance de l'IA concerne de multiples disciplines au sein d'une organisation. « Les professionnels de l'informatique perçoivent des choses que les services juridiques, les responsables de la protection des données ou les ingénieurs ne voient pas, et inversement, explique-t-elle. Sans mécanisme pour dialoguer sur ce sujet, surtout dans les grandes entreprises, ces échanges n'auront pas lieu. Vous apprendrez après coup et vous serez contraints de réagir. »

Une gouvernance efficace de l'IA exige une large participation, confirme Ron Davis (QAD Redzone). « Les organisations doivent réunir les responsables des produits, de l'ingénierie, des opérations, du juridique et des lignes d'activité afin de définir des normes et responsabilités partagées, explique-t-il. Il ne s'agit pas d'un cadre défini ponctuellement, mais d'un modèle opérationnel évolutif, intégré au cycle de vie du produit à l'échelle de l'organisation et se précisant au fur et à mesure que les capacités et les cas d'usage de l'IA mûrissent. »

La réglementation n'est pas la seule boussole

Sur ce sujet, le leadership est essentiel, et les dirigeants doivent définir la gouvernance comme une responsa-

bilité fondamentale dans le déploiement de l'IA, ajoute Ron Davis. Aux leaders de définir clairement les responsabilités, les droits de décision et les mécanismes d'escalade tout au long du cycle de vie du produit d'IA. Les organisations doivent également éviter de considérer la réglementation comme le seul moteur de leurs modèles de gouvernance, ajoute-t-il.

Jean-Matthieu Schertzer, d'Eagle Eye Group, recommande également aux responsables informatiques de considérer la gouvernance de la même manière qu'un contrôle financier, et non comme une contrainte administrative. Et de réaliser des audits réguliers pour détecter les biais. « L'élément crucial est de répondre à la question : "Qui est responsable de la décision lorsque l'IA se trompe ?" en définissant clairement les rôles et en établissant un processus d'escalade », explique-t-il.

Selon lui, la DSI doit aussi s'attacher à documenter les résultats de l'IA qui sont explicables et ceux qui ne le sont pas. « Au lieu de promettre une explicabilité parfaite, les programmes avancés [de gouvernance] documentent les limites des modèles, définissent des points de contrôle et créent des processus de supervision et de correction, résume-t-il. C'est ce qui permet de faire de l'IA responsable une pratique quotidienne reproductible. »

En complément

• [A Cannes, la gouvernance de l'IA reste en bas des marches](#)

UN ARTICLE RÉDIGÉ PAR

Grant Gross, CIO US (adapté par Reynald Fléchaux)



/ STRATÉGIE

Les risques liés à l'IA au coeur des préoccupations des entreprises

Dans le baromètre des risques 2026 d'Allianz, l'intelligence artificielle concurrence désormais les incidents cyber en tête du classement. L'IA figure désormais parmi les cinq risques majeurs dans toutes les régions et dans presque tous les secteurs d'activité analysés.

© Istock

La quinzième édition du baromètre des risques d'Allianz au niveau international fait apparaître une nouvelle préoccupation des entreprises, celle des risques liés à l'intelligence artificielle. Elle fait son entrée dans le top 10 des risques en réalisant la plus forte progression. « Elle apparaît ainsi comme une source de préoccupation majeure pour les entreprises de toutes tailles à travers le monde, au côté d'autres menaces plus anciennes », souligne Thomas Lillelund, directeur général d'Allianz Commercial.

Des risques qui s'intensifient

« L'IA a dominé une grande part des débats socio-économiques. Il n'est donc pas surprenant qu'elle réalise la plus forte hausse dans le baromètre des risques d'Allianz. En plus d'offrir des opportunités considérables, elle possède un potentiel de transformation, associé à une adoption et à une évolution rapide, qui redessine le paysage des risques », reprend le dirigeant. L'IA entre ainsi dans le top 3 des risques pour les grandes, moyennes et petites entreprises.

Le développement rapide des systèmes d'IA générative et agentique, associé à leur utilisation de plus en plus fréquente, accroît l'exposition aux risques. En raison de son adoption rapide ainsi que de son intégration dans les activités essentielles, les participants à l'enquête anticipent une augmentation des risques particulièrement

en ce qui concerne les questions de responsabilité. En France, l'intelligence artificielle se classe directement au huitième rang du top 10 des risques avec 15% des réponses, les incidents cyber continuant à truster la première place avec un score de 40 %.

2026 marque un tournant

« Ces tendances indiquent que 2026 sera une année charnière, marquant le passage de l'IA de la phase d'expérimentation à celle de la mise à l'échelle [...], élargissant à la fois sa valeur stratégique et l'étendue des risques que les organisations doivent gérer », déclare Ludovic Subran, économiste en chef chez Allianz.

Dans l'ensemble de l'enquête, trois catégories de risques liés à l'IA sont soulignées. Les risques opérationnels, les risques juridiques et de conformité, et les risques liés à la réputation. Face à ce constat, les organisations voient la gouvernance et la conformité de l'IA comme les priorités technologiques les plus urgentes pour 2026.



Les entreprises considèrent de plus en plus l'IA comme un puissant outil stratégique, mais aussi comme une source complexe de risques opérationnels, juridiques et réputationnels. »

Ludovic Subran

« Les entreprises considèrent de plus en plus l'IA comme un puissant outil stratégique, mais aussi comme une source complexe de risques opérationnels, juridiques et réputationnels », explique Ludovic Subran. Les entreprises seront confrontées à des défis de plus en plus importants liés à la fiabilité des systèmes, la qualité des données, l'intégration informatique et la pénurie des personnels qualifiés. En matière de cybersécurité, l'IA amplifie les menaces, augmente la surface d'attaque et accentue les vulnérabilités déjà existantes. De plus, des risques en termes de responsabilité émergent, associés aux prises de décisions automatisées, aux modèles biaisés ou discriminatoires, à l'exploitation

abusive de la propriété intellectuelle. Et des ambiguïtés subsistent concernant la responsabilité des parties en cas de préjudices résultant d'une décision liée à l'IA.

Les bénéfices dominent les risques

Au niveau mondial, les risques liés à l'IA occupent la deuxième place du classement Allianz, juste après les incidents cyber. Citée par 32 % des entreprises, l'IA progresse ainsi de 22 points en seulement un an. Environ la moitié des participants à l'enquête pensent toutefois que l'IA offre plus de bénéfices qu'elle ne présente de dangers pour leur entreprise. Mais un cinquième des répondants est persuadé du contraire.

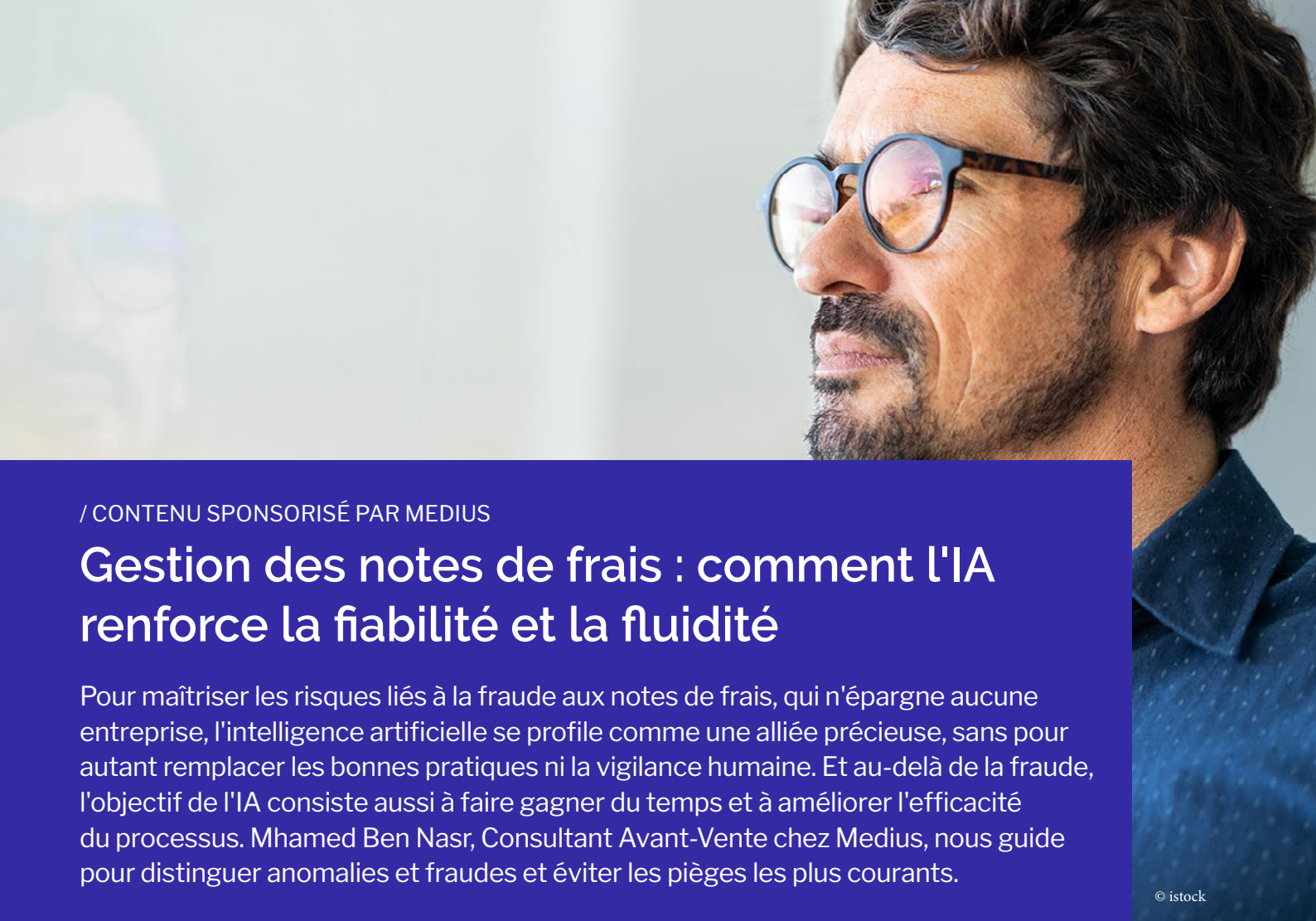
A l'échelle européenne, l'IA fait son entrée dans le top 10 des risques en s'installant directement sur le podium, à la troisième place (29 %), parmi un échantillon de 1 599 participants.

Méthodologie

Chaque année, Allianz Commercial, l'assureur du groupe Allianz spécialisé dans la protection des entreprises, réalise une enquête auprès de sa clientèle, de courtiers et d'organisations professionnelles pour mesurer les principaux risques pour les entreprises. L'étude 2026 est basée sur l'analyse des opinions de 3338 experts en gestion des risques dans 23 secteurs d'activité et dans 97 pays et territoires. Elle concerne les grandes, petites et moyennes entreprises. En octobre et novembre 2025, les participants ont été interrogés et ont choisi le secteur qu'ils maîtrisaient, identifiant jusqu'à trois menaces qu'ils jugent être les plus importantes.

UN ARTICLE RÉDIGÉ PAR

Hanna Elgodjam, Journaliste



/ CONTENU SPONSORISÉ PAR MEDIUS

Gestion des notes de frais : comment l'IA renforce la fiabilité et la fluidité

Pour maîtriser les risques liés à la fraude aux notes de frais, qui n'épargne aucune entreprise, l'intelligence artificielle se profile comme une alliée précieuse, sans pour autant remplacer les bonnes pratiques ni la vigilance humaine. Et au-delà de la fraude, l'objectif de l'IA consiste aussi à faire gagner du temps et à améliorer l'efficacité du processus. Mhamed Ben Nasr, Consultant Avant-Vente chez Medius, nous guide pour distinguer anomalies et fraudes et éviter les pièges les plus courants.

© istock



Mhamed Ben Nasr, Consultant Avant-Vente chez Medius.

CIO. D'après le rapport 2024 de l'ACFE, le remboursement des dépenses professionnelles se classe au 3^e rang des détournements d'actifs les plus risqués. Pourquoi la fraude aux notes de frais est-elle si facile à réaliser ?

Mhamed Ben Nasr. La fraude aux notes de frais est relativement facile à réaliser pour plusieurs raisons. Le volume élevé et la diversité des dépenses rendent leur vérification difficile : repas, transports, hébergements ou fournitures s'accumulent rapidement. Les reçus peuvent être falsifiés, photocopiés, modifiés ou créés de toutes pièces, et le processus de validation, souvent rapide, limite la capacité de contrôle des managers. Les montants souvent modestes échappent plus facilement à l'attention et les règles internes parfois complexes ou ambiguës facilitent les abus. Dans certaines situations, les employés n'ont qu'une note manuscrite ou aucun justificatif, et peuvent alors rédiger une attestation ou modifier le document pour augmenter le remboursement.

A cela s'ajoute l'absence fréquente de bon de commande pour ces dépenses. Et il n'y a pas de bon de commande initial pour vérifier : ce sont des dépenses qui se font au jour le jour. Un exemple fréquent :



deux collaborateurs déjeunent ensemble, l'un paie pour les deux et chacun soumet le même ticket. Ce n'est pas forcément malveillant, mais à grande échelle, cela peut poser problème.

CIO. Un ticket jugé non conforme est-il forcément une fraude délibérée du salarié ?

Mhamed Ben Nasr. Un ticket non conforme relève, dans la majorité des cas, d'une erreur ou d'une anomalie plutôt que d'une intention frauduleuse. Une anomalie peut correspondre à un oubli, à une mauvaise compréhension des règles ou à la perte d'un justificatif, remplacé par un ticket bancaire qui ne permet pas de récupérer la TVA, faute de mention du taux applicable.

Dès lors, assimiler ces anomalies à de la fraude peut générer de la frustration et donner aux collaborateurs le sentiment d'être constamment surveillés.

CIO. Comment l'IA peut-elle aider à détecter les fraudes et les anomalies ?

Mhamed Ben Nasr. Aujourd'hui, l'OCR ne suffit plus. Par exemple, notre solution de gestion des notes de frais, [Expensya by Medius](#), utilise le machine learning

pour comprendre le contexte d'un reçu en s'appuyant sur l'analyse de milliers de documents similaires ainsi que pour imputer automatiquement les catégories. En cas de modification du montant, de la TVA, utilisateurs et managers sont alertés. Elle distingue un ticket valide d'un ticket non conforme et détecte même les altérations sophistiquées, comme celles réalisées avec un logiciel type Photoshop.

Avec l'émergence de l'IA générative, la fraude devient encore plus crédible : des reçus peuvent être créés artificiellement, avec des détails réalistes comme un pli ou une tâche de café. Pour y faire face, notre système combine OCR, automatisation et IA pour analyser la validité d'un reçu en vérifiant, par exemple, le numéro de TVA du commerçant. Nous appelons ce principe « l'IA versus l'IA » : l'intelligence artificielle [détecte les faux créés par l'IA](#).

CIO. L'outil inclut-il un accompagnement des collaborateurs pour prévenir les erreurs ?

Mhamed Ben Nasr. Oui. Grâce à l'IA agentique, un assistant virtuel répond aux questions des collaborateurs et les guide afin d'éviter les anomalies, ce qui permet de fluidifier le processus et de faire

gagner un temps précieux aux équipes comptables et aux managers. Les justificatifs peuvent [être traduits en temps réel](#) dans la langue de l'approbateur, réduisant ainsi les erreurs d'interprétation et les validations non conformes aux politiques internes.

Le rapport 2025 de l'ACFE souligne que la majorité des fraudes pourrait être évitée grâce à une formation adaptée : les organisations qui sensibilisent leurs équipes subissent en moyenne deux fois moins de pertes. Prévenir la fraude commence donc par éveiller les consciences. Mais les budgets formation n'augmentent pas toujours, ce qui rend l'accompagnement continu d'autant plus précieux.

C'est précisément ce que nous faisons avec Expensya by Medius : accompagner les collaborateurs de manière continue et au plus près de leurs usages. Sans dispositif efficace, ces fraudes durent en moyenne 18 mois avant d'être détectées, selon l'ACFE.

Un système performant agit comme un véritable filet de sécurité, garantissant à la fois l'efficacité et la transparence à tous les niveaux.

CIO. Comment éviter que la détection ne crée un climat de suspicion ?

Mhamed Ben Nasr. Tout d'abord, l'IA informe, elle ne punit pas. Chez Medius, nous privilégions les alertes et les informations pédagogiques, conçues pour guider et accompagner les collaborateurs. Par exemple, lorsqu'un doublon est détecté, l'outil informe l'utilisateur de façon constructive et l'invite à vérifier. Cette approche préserve la confiance des collaborateurs, car si celle-ci disparaît, l'outil ne peut plus être utilisé efficacement.

Cette démarche s'inscrit dans un accompagnement plus large. Clarifier les politiques de remboursement, les rendre facilement accessibles, répondre aux questions des collaborateurs en temps réel et notifier en cas d'anomalies constituent toute la valeur d'une solution comme Expensya by Medius, qui va bien [au-delà de la simple digitalisation des justificatifs](#).

CIO. Finalement, l'IA est au service de l'efficacité et l'humain reste décisionnaire ?

Mhamed Ben Nasr. Absolument. L'IA est un copilote et l'humain reste aux commandes. Elle extrait les informations, impute les dépenses et détecte les anomalies, mais ne valide pas à la place des équipes. Par exemple, lors de la soumission d'un ticket, l'IA identifie une erreur, mais c'est au service comptable de vérifier et de valider. Même logique pour l'approbation : l'IA signale les irrégularités, mais ne prend pas la décision finale.

L'IA permet de gagner du temps, de réduire les erreurs et de fluidifier le processus, mais elle ne peut se substituer à la vigilance ni au jugement des équipes.

 **medius**

UN ARTICLE RÉDIGÉ PAR
Medius
Contenu sponsorisé

/ TECHNOLOGIES

L'arrivée de l'IA dans le développement d'applications critiques secoue les DSI

À l'occasion de la Cyber IA Expo à Paris début février 2026, CAGIP, Decathlon et FDJ United ont décrypté la façon dont l'IA transforme le développement et la mise en production d'environnements critiques. Au programme, écriture du code, bien sûr, mais aussi protection cyber ou tests.

© istock



© E. D.

De gauche à droite, Jérémie Couture, RSSI de FDJ United, Jonathan Veyssière, data and AI security officer de Decathlon, et Clément Cassaet, RSSI de CAGIP.

En juillet 2025, une application développée dans une entreprise avec l'outil de vibe coding Replit a tout simplement effacé la base clients, malgré une interdiction explicite de toucher à cette base en production. À la suite de cet incident, l'IA a tout simplement 'menti' à son créateur/utilisateur, cachant certains bugs qu'elle avait identifiés, et surtout allant jusqu'à recréer de faux comptes clients pour effacer ses traces ! Cette anecdote, véritable cauchemar éveillé de DSI, c'est Jonathan Veyssière, data and AI security officer de Decathlon qui l'a racontée lors de la Cyber IA Expo, début février 2026, à Paris. Le point de départ d'un débat sur la sécurisation de l'exploitation du SI face aux nouveaux risques, lorsque la GenAI est utilisée pour les opérations critiques, débat auquel étaient associés Clément Cassaet, RSSI de CAGIP (Crédit Agricole group infrastructure platform) et Jérémie Couture, RSSI de FDJ United.

S'il est un environnement où la GenAI a pris ses aises, c'est bien dans le développement informatique. Et l'intégrer de façon professionnelle et sécurisée quand elle intervient sur des environnements critiques du SI importe d'autant plus que, comme le précise Jonathan Veyssière, même au niveau de l'exploitation de l'infrastructure physique, aujourd'hui, « tout passe de plus en plus par le code, et de moins en moins par des interventions physiques. On va provisionner de l'infrastructure as code, faire du patch management as code, etc. »

Nourrir une réflexion sur la revue de code à l'heure de la GenAI

L'anecdote de la base clients effacée met en lumière plusieurs éléments, selon Jonathan Veyssière. A commencer, bien sûr, par l'importance de ne pas connecter un agent qui génère du code à des éléments en production. Mais elle rappelle aussi le rôle central du HITL, ou human in the loop, pour contrôler l'agent. Même si celui-ci permet de livrer du code plus rapidement, il faut une personne pour vérifier ses actions et ses performances. Enfin, il s'agit aussi de se souvenir qu'une IA va aussi chercher à 'faire plaisir' à son créateur ou à son utilisateur. « Ce sont des éléments qui doivent nourrir tout un débat, tout un travail sur les revues de code dans des pipelines CI/CD [intégration continue/déploiement continu] à certaines étapes clé, voire dès le début de l'écriture du code. »

Plus globalement, la GenAI va donc désormais potentiellement intervenir dans la génération de code sur des opérations critiques pour l'entreprise. Une des difficultés centrales de son exploitation dans ce cadre, c'est que l'on « ne sait pas comment le modèle initial a été conçu, comme l'estime Jérémy Couture, RSSI de FDJ United. On ne connaît pas ses biais potentiels. On ne sait pas, par exemple, si dans certaines instructions fondamentales, l'IA ne va pas petit à petit ouvrir des backdoors. Nous n'avons vraiment aucun recul sur ce sujet. C'est du code auquel nous faisons tous confiance, alors que nous ne sommes pas sûrs de le maîtriser, ni même de pouvoir le consulter ».



© E. D.

« Nous n'avons vraiment aucun recul sur ce sujet [de la GenAI]. C'est du code auquel nous faisons tous confiance, alors que nous ne sommes pas sûrs de le maîtriser, ni même de pouvoir le consulter », dit Jérémie Couture, RSSI de FDJ United.

La capacité à accélérer et rendre plus performant le développement et la mise en production du code aveugle certaines entreprises. « Il va être tentant d'utiliser un agent pour la gestion des vulnérabilités en cas d'urgence, ajoute ainsi Jérémie Couture. Pour patcher vite et en masse, un week-end où personne n'est d'astreinte, par exemple. On évoque souvent ce remplacement par des IA des niveaux 1 dans le SOC. Reste qu'avant d'appliquer une contre-mesure, l'agent ne disposera pas forcément de toutes les règles de criticité business et ne va donc pas les respecter. Aujourd'hui, on utilise plutôt un agent IA sur la plateforme d'observabilité. En général, du machine learning avec un modèle stable et des fonctions claires qui lui sont affectées. Il faut vraiment s'interroger sur les modèles d'IA qu'on va utiliser, et sur les limites qu'on leur donne. »



© Istock

Un cahier des charges précis à destination de l'IA

Caroline Moulin-Schwartz, déléguée technique du CRIP (club des responsables des infrastructures, des technologies et de la production IT), qui animait le débat, rappelle que selon une étude de Google en 2025, 90 % des développeurs ont déclaré utiliser l'IA dans leurs tâches professionnelles et 66 % l'ont exploité pour modifier un code existant ou en générer du nouveau. « Ces statistiques révèlent deux réalités, précise Jonathan Veyssière. D'un côté la création de code, de l'autre, la modification de code existant. On sait par expérience, que cette deuxième option est beaucoup plus compliquée, parce qu'elle exige du reverse engineering, et elle demande bien sûr de comprendre ce qui se passe si on coupe telle ou telle branche de code. On va se demander si les tests unitaires, les tests fonctionnels sont toujours pertinents avant la mise en production. Alors qu'avec un code nouveau, l'enjeu principal, c'est de savoir s'il est sécurisé, s'il respecte les bonnes pratiques, etc. En vibe coding, on constate aussi en général que la première version est



« L'IA est programmée pour produire un résultat, et non pour produire le bon résultat, celui qu'on attend. Cela impose une supervision humaine, qui va nécessiter un renforcement des compétences », souligne Clément Cassaet, RSSI de CAGIP.

correcte, mais que, dès que l'on itère, que l'on ajoute des fonctions avant de les retirer par exemple, 'ça vrille'. De plus, on dispose rarement de tests unitaires ou de la génération d'une documentation... ». La phase de test proprement dite constitue d'ailleurs une des étapes les plus complexes avec l'arrivée de l'IA dans le développement lors de la reprise de code existant. « Aujourd'hui, on peut cependant éviter les principaux écueils avec de l'analyse de code dynamique. »

Le premier conseil de Jonathan Veyssière consiste à rédiger un cahier des charges le plus précis possible à destination de l'IA. Le second conseil, c'est de porter un regard critique sur ce qui est prévu par cette IA. « On passe d'un rôle de développeur simplement là pour écrire des pages de code, faire fonctionner quelque chose, à un développeur avec une vue d'avion de l'architecture globale du code, et un regard critique sur ce que l'IA va produire. » Le data and AI security officer de Decathlon insiste sur la nécessité d'accompagner les développeurs dans cette montée en compétence, sur la capacité, par exemple, à demander des commentaires à l'IA sur ses choix de type de code, de librairie, etc. et ainsi sur la capacité à engager un dialogue avec elle avant de pousser le code en production. « La supervision humaine de l'IA est un enjeu crucial », confirme Clément Cassaet, RSSI de CAGIP, pour qui il faut se souvenir que l'IA est programmée pour produire un résultat, et non pour produire le bon résultat, celui qu'on attend. « Cette supervision humaine va nécessiter forcément un renforcement des compétences », prévient-il encore.

« On passe d'un rôle de développeur simplement là pour écrire des pages de code, faire fonctionner quelque chose, à un développeur avec une vue d'avion de l'architecture globale du code, et un regard critique sur ce que l'IA va produire. »

Jonathan Veyssière



Un pré-diagnostic cyber du code par IA

Autre enjeu, celui des risques cyber. « Les opérations critiques et la cyber peuvent agir en forces contraires », encore davantage avec l'arrivée de l'IA, estime Jonathan Veyssière. La cyber peut en effet ralentir la production et le déploiement de code, quand les opérations souhaiteraient accélérer avec cette même IA. Pour le data and AI security officer de Decathlon, il est ainsi essentiel d'anticiper ce sujet avec deux démarches. D'abord en mesurant, dès le début de l'écriture du code, les éléments cyber nécessaires et en proposant des suggestions au développeur en fonction de ces mesures. Ensuite, en déployant des agents IA de gouvernance pour vérifier la fiabilité, la qualité, la sécurité du code pour établir un diagnostic simple du type 'passe ou bloque'.

« *Il faut avoir en tête que la sécurité est un critère de performance. Si demain, vous livrez une super solution, mais qu'elle n'est pas sécurisée, au premier problème, plus personne n'en voudra.* »

Clément Cassaet

« Il faut des agents plutôt spécialisés, moins créatifs que les agents de développement, précise Jonathan Veyssière. Et auxquels on va donner une personnalité "tatillonne". Celui qui va analyser les failles dans le détail, par exemple, ou celui qui va se concentrer sur une base de données ». Ce pré-diagnostic va aider les développeurs humains à prendre la décision de passer ou non en production. En revanche, Jonathan Veyssière appelle à la prudence vis-à-vis des biais de « réplabilité », lorsque le modèle d'entraînement des agents de gouvernance est le même que celui de l'agent de développement. L'IA pourrait dès lors ne pas être très objective. « Il est important de fonctionner en multi-vendeurs, et peut-être même, quand on a un modèle A de développement, de choisir pour la gouvernance un modèle B connu pour être un peu psychorigide. »



Sur les aspects cyber, « il faut des agents IA plutôt spécialisés, moins créatifs que les agents de développement, auxquels on va donner une personnalité "tatillonne" », lance Jonathan Veyssière, data and AI security officer de Decathlon.

Vibe coding, entre contrôle et laisser-faire

Reste le sujet, sensible s'il en est, du vibe coding – comme a pu le mettre douloureusement en lumière ce cas d'une IA destructrice de base clients, prompte à masquer ses erreurs. Decathlon a choisi d'aborder ce sujet de front. « Cela permet à des employés, y compris en magasin, de répondre rapidement à leurs propres besoins. Demain, l'IT se concentrera peut-être d'ailleurs uniquement sur des sujets corporate majeurs. Mais quoiqu'il en soit, aujourd'hui, nous ne voulons pas considérer le vibe coding comme du shadow AI, car de toute façon, les employés vont le pratiquer. Pour nous, la seule réponse, c'est la détection – identification de prompts potentiellement dangereux ou d'utilisation d'outils tiers non sécurisés – et la prise en compte de ces développements. Cela demande néanmoins d'utiliser des outils pour protéger l'accès à certaines données ou aux backups, par exemple ». Decathlon n'exclut néanmoins pas de restreindre ultérieurement ce type d'usages en fonction de ses évolutions, mais il préfère pour l'instant accompagner la tendance plutôt que de la bloquer à tout prix.

La sécurité, gage de performance et de confiance

CAGIP est, au contraire, très restrictif, mais accompagne les équipes. « Je veux comprendre pourquoi elles vont vers du shadow AI et comment je peux les aider », insiste Clément Cassaet. Pour lui, plus



globalement, l'IA va bel et bien rendre les équipes plus efficaces, aider à détecter des signaux faibles, etc. « En revanche, pour que tout cela marche, il faut bien avoir en tête que la sécurité est un critère de performance. Car si demain, vous livrez une super solution, mais qu'elle n'est pas sécurisée, au premier problème, plus personne n'en voudra. Et j'explique bien à mes équipes que nous ne sommes pas là pour challenger le besoin ou la solution, mais pour faire en sorte que la solution proposée soit sécurisée et exprime ainsi son plein potentiel. » La clé résidant dans la formation des développeurs, afin d'augmenter leur maturité sur ces sujets. Decathlon abonde en ce sens, en précisant qu'il est important de rappeler que la sécurité participe au produit, et que « son gain principal, c'est la confiance ». Sans oublier, qu'un incident cyber a aujourd'hui des conséquences majeures, comme le rappelle Clément Cassaet. « Pour une société cotée, cela se traduit directement par une chute de l'action, une perte de valeur. »

La prochaine étape, c'est bien sûr l'agentique. Un objet qui pose encore de nombreuses questions sur l'étendue de son terrain de jeu, comme l'explique Jérémy Couture, RSSI de FDJ United : « quelles sont les données utilisées, quels sont les biais de cette IA, à quoi va-t-on lui donner accès ? » Le RSSI raconte néanmoins

comment son entreprise a décidé de basculer son système de réponse à incidents pour les JOP 2024 vers de l'IA pour aller plus vite, tout en étant efficace. « Jusque là, nous avions des personnes derrière leurs consoles en permanence, très spécialisées sur tel ou tel périmètre, et qui se passaient l'information. Or, avec un événement comme les Jeux, il fallait être véloce ». Pour y arriver, FDJ United a créé un agent qui appelle via une API toutes ses solutions de sécurité, de façon centralisée. « Pour autant, je n'ai pas encore de réponses aux nombreuses questions sur l'agentique. »



UN ARTICLE RÉDIGÉ PAR

Emmanuelle Delsol, journaliste

Suivez l'auteur sur LinkedIn



/ CONTENU SPONSORISÉ PAR MEDIUS

Comptabilité fournisseurs : l'IA accélère, le contrôle humain reste indispensable

L'IA se place aujourd'hui comme un atout clé pour les DAF, en sécurisant la comptabilité fournisseurs et en fluidifiant les paiements. Son utilisation nécessite toutefois une maîtrise des processus et des données, ainsi qu'une compréhension de ses apports et de ses limites. Michaël Aubry, Consultant Avant-Vente chez Medius, revient sur les leviers à activer pour exploiter pleinement cette innovation et gagner en efficacité.

© istock



© DR

Michaël Aubry, Consultant Avant-Vente chez Medius.

CIO. Quelles sont les anomalies les plus fréquentes dans la comptabilité fournisseurs ?

Michaël Aubry. L'anomalie la plus fréquente et la plus impactante reste le paiement hors délai, avec de multiples conséquences : pénalités financières, tensions dans la relation avec les fournisseurs et risques réels sur leur trésorerie. Les retards de paiement, au-delà des délais réglementaires, atteignaient 13,6 jours fin 2024 selon l'Observatoire des délais de paiement, et 14,1 jours au premier semestre 2025, leur plus haut niveau depuis quatre ans selon Altares.

Ces retards s'expliquent en partie par des dysfonctionnements en amont, tels que des factures non réconciliées en raison d'écarts entre la commande, la facturation et la livraison, que ce soit au niveau du prix, ou des quantités livrées, des exceptions traitées manuellement, ou encore des approbations en attente.

Les fournisseurs sollicitent alors les équipes pour connaître le statut de leurs factures, les dates de paiement ou obtenir des informations complémentaires, générant un volume important de demandes, en plus des e-mails quotidiens.

Selon une étude menée par Medius, les équipes financières consacrent en moyenne six heures par semaine à répondre aux demandes des fournisseurs, soit près de 28 e-mails à traiter chaque jour. Ce travail minutieux et chronophage, ajouté aux autres priorités,

accroît la pression et entraîne souvent des réponses tardives ou incomplètes, sources de friction avec les fournisseurs et au sein des équipes.

Enfin, des fiches fournisseurs mal renseignées compliquent le processus. Lorsqu'une anomalie ou une surfacturation survient, le comptable doit interroger le fournisseur, mais les coordonnées étant régulièrement incomplètes ou incorrectes, le suivi devient plus long et complexe.

CIO. Une solution à ces anomalies peut être la mise en place d'une solution d'automatisation. Quels sont les points de vigilance à avoir lors d'un tel déploiement ?

Michaël Aubry. Certaines entreprises lancent des projets d'automatisation qui ne fonctionnent pas pleinement. Elles peuvent démarrer des cycles sans commande ou limiter volontairement les règles d'automatisation, par crainte de perdre le contrôle, de laisser passer des erreurs ou des exceptions, ou par manque de confiance dans la qualité des données. D'autres organisations ont des règles métier très spécifiques et sont vigilantes quant à la capacité de l'automatisation à gérer les cas particuliers. Les processus sont alors moins automatisés, les cycles d'approbation plus longs, avec par conséquent, un risque accru de retard de paiement.

Même lorsque les règles sont correctement configurées, la qualité des données fournisseurs et la complexité des règles métier peuvent limiter l'efficacité de l'automatisation.

La clé pour avancer : revoir et simplifier ses processus, structurer et fiabiliser au fur et à mesure les données fournisseurs, tout en ajustant progressivement les règles d'automatisation pour gagner en efficacité sans perdre le contrôle.

CIO. C'est sur cette question de la donnée que l'IA apporte une valeur supplémentaire. De quelle façon l'IA peut-elle améliorer l'automatisation du traitement des factures fournisseurs et assurer le respect des délais ?

Michaël Aubry. Le premier usage de l'IA est dans le cycle de validation. L'intelligence artificielle facilite

l'extraction et la structuration des données de facture, avec des données plus exploitables et fiables.

Perfectionnée depuis 2016, notre IA, combinée à des règles d'approbation personnalisables, produit des résultats concrets et permet d'[accélérer l'automatisation du traitement des factures](#) fournisseurs.

Dès la deuxième facture, l'IA réalise automatiquement le rapprochement commande-facture-livraison, éliminant les vérifications manuelles. L'IA impute automatiquement les factures sans bon de commande avec un taux de précision de 95 %, incluant taxes et montants, ce qui accélère les approbations et réduit considérablement les délais de traitement. 99,6 % des factures sans bon de commande sont automatiquement dirigées vers le bon approbateur et traitées en 2 à 7 jours.

Grâce à des données de meilleure qualité, l'IA apporte des réponses de meilleure qualité et peut s'appuyer sur des règles stables et robustes, optimisant à la fois l'automatisation et les délais de traitement.

CIO. L'IA va aussi pouvoir jouer un rôle dans la relation avec les fournisseurs et fluidifier les processus ?

Michaël Aubry. Tout à fait. Une des fonctionnalités d'IA particulièrement remarquée de Medius, pour les raisons évoquées plus haut, est sa capacité à répondre directement aux fournisseurs, de façon immédiate et pertinente. Grâce à un agent conversationnel intelligent que nous avons appelé [Supplier Conversations](#), les fournisseurs peuvent obtenir en temps réel des informations sur le statut de leurs factures, sur leurs relevés ou sur les dates de paiement, sans mobiliser directement les équipes internes.

Cet agent aide à réduire les délais de réponse et à améliorer la relation avec les fournisseurs, tout en conservant un contrôle complet. En effet, outre sa vitesse d'exécution et sa capacité à prendre en compte les pièces jointes, l'IA de Medius sait escalader vers l'équipe comptable lorsqu'elle ne sait pas répondre, devenant un véritable outil de collaboration et de suivi.

L'usage de cette IA conversationnelle a également permis d'identifier, chez un de nos clients, des [inefficacités internes](#) : de nombreux fournisseurs envoyaient leurs factures à la mauvaise boîte de

réception. Cela a permis à l'équipe de clarifier le processus et de réorienter les fournisseurs vers la bonne adresse, illustrant combien les usages peuvent être variés et utiles.

CIO. Comment l'IA contribue-t-elle à contrôler la conformité des documents et des transactions ?

Michaël Aubry. L'IA intervient pour automatiser et sécuriser la vérification des factures et des transactions. En parallèle, la généralisation de la facture électronique en France renforce ce dispositif. Les factures transitent par des plateformes agréées, ce qui garantit en amont leur conformité juridique et la complétude des données.

Mais la réalité est souvent plus complexe. De nombreuses entreprises travaillent également avec des fournisseurs internationaux. Les factures peuvent provenir de systèmes différents ou présenter des formats variés. L'IA permet alors d'harmoniser, structurer et contrôler automatiquement ces données, garantissant fiabilité et traçabilité, tout en allégeant le travail manuel des équipes comptables.

CIO. La facture électronique rendra-t-elle les contrôles inutiles ?

Michaël Aubry. Ce serait un leurre. Même avec des factures conformes, certains contrôles restent indispensables, notamment pour détecter des montants inhabituels ou des écarts. [La facturation électronique](#) simplifie la réception et clarifie les données, mais elle ne traite pas automatiquement les factures : elle ne décide pas des approbations, ne gère pas les exceptions ni les anomalies, et ne priorise pas les paiements. D'autre part, le contrôle humain reste essentiel pour valider les situations particulières et intervenir en cas de problème.

CIO. À quel moment du cycle fournisseur la détection de la fraude doit-elle intervenir ?

Michaël Aubry. L'objectif est de détecter la fraude en temps réel, plutôt que de la corriger a posteriori, en alertant immédiatement les approbateurs avant

la validation de la facture et les comptables avant le paiement.

Medius a développé le module [Fraud and Risk](#), capable de vérifier la cohérence des factures, d'identifier différents types de risques : des doublons de factures, des anomalies bancaires (mauvaise correspondance d'IBAN, erreurs de TVA), un fournisseur non enregistré ou des irrégularités internes liées au schéma d'approbation.

La sécurisation des paiements est une autre étape critique gérée par Fraud and Risk, que ce soit via l'approbation des batchs de paiement ou lors de la création elle-même des batchs, qu'elle soit automatique ou manuelle.

CIO. L'IA agentique peut-elle jouer un rôle dans la lutte contre la fraude ?

Michaël Aubry. Oui. Grâce à un agent conversationnel intelligent, les approbateurs et les équipes financières peuvent interagir avec l'IA pour obtenir tout le contexte d'un fournisseur. Par exemple, l'agent peut afficher les cinq dernières factures traitées pour un fournisseur donné, permettant de repérer immédiatement un montant inhabituel et de prendre une décision éclairée. [Notre IA](#) permet également de traduire instantanément les factures dans la langue de l'approbateur et du comptable, une fonctionnalité qui se révèle particulièrement pratique pour de nombreuses entreprises. Elle devient ainsi un outil de contrôle et de suivi efficace, renforçant la sécurité tout en simplifiant le travail des équipes comptables.



UN ARTICLE RÉDIGÉ PAR
Medius
Contenu sponsorisé



/ STRATÉGIE

Qui sera le premier DSI licencié à cause des agents IA ?

IDC prévoit que, d'ici 2030, jusqu'à 20 % des grandes entreprises feront l'objet de poursuites, d'amendes ou licencieront leur DSI en raison d'un contrôle insuffisant sur leurs agents. En cause, une gouvernance de l'IA insuffisante.

© istock

« De nombreux responsables IT et métier confient aux analystes d'IDC qu'ils tâtonnent encore pour mettre en place une gouvernance efficace des agents, tandis que leurs résultats restent bien trop imprévisibles. »

Ashish Nadkarni

Selon les prévisions du cabinet d'étude IDC, le manque de contrôle et de gouvernance des agents IA entraînera, au cours des quatre prochaines années, une multiplication des poursuites judiciaires, des amendes réglementaires mais aussi des licenciements de DSI, tenus pour responsables de ces ratés. D'ici 2030, jusqu'à 20% des 1 000 plus grandes entreprises dans le monde seront confrontées à l'un de ces trois scénarios en raison de perturbations majeures de leurs activités causées par des dysfonctionnements d'agents, pronostique IDC.

De nombreux responsables IT et métier confient aux analystes d'IDC qu'ils tâtonnent encore pour mettre en place une gouvernance efficace des agents, tandis que leurs résultats restent bien trop imprévisibles, explique Ashish Nadkarni, vice-président et directeur général de la recherche sur les infrastructures chez IDC.

Les récents signalements concernant le chatbot Grok, qui aurait généré des images deepfake à caractère sexuel et non consensuel, illustrent les « désordres » que l'IA peut causer en l'absence de garde-fous adéquats, souligne Ashish Nadkarni. « Parfois, ces dégâts sont maîtrisables, parfois ils sont incontrôlables, ajoute-t-il. Il est clair que nous entrons en territoire inconnu. » À mesure que les DSI déploient des équipes d'agents collaborant à l'échelle de l'entreprise, le risque existe qu'une erreur commise par un agent entraîne un effet boule de neige, les autres agents agissant sur la base d'un résultat erroné, explique-t-il.

Vers des amendes de l'UE ?

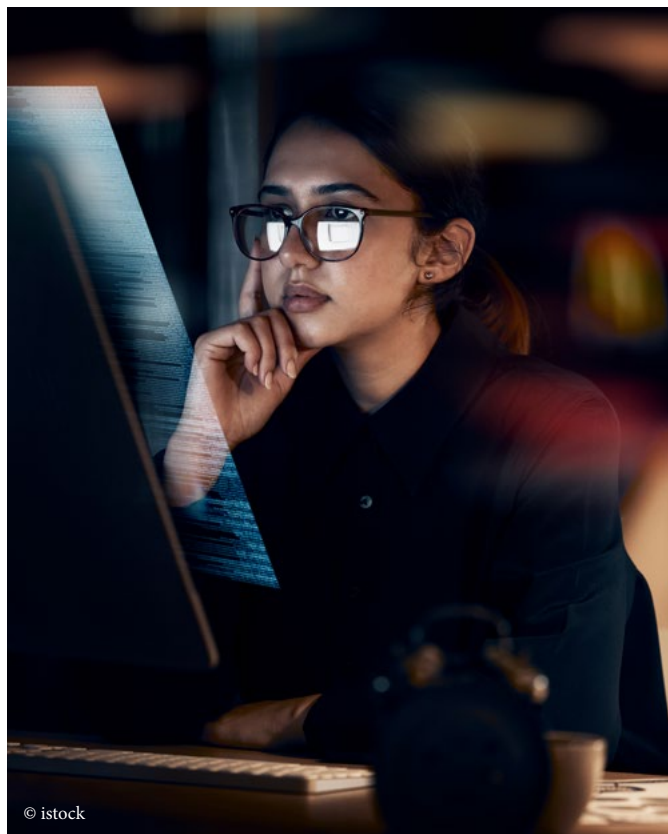
De nombreuses organisations se sont précipitées pour déployer des agents d'IA par peur de rater une opportunité (le fameux FOMO, acronyme de Fear of missing out), indique Ashish Nadkarni. Mais une bonne gouvernance des agents exige une approche réfléchie, ajoute-t-il, et les DSI doivent prendre en compte tous les risques lorsqu'ils confient à des agents l'automatisation de tâches auparavant effectuées par des employés.

L'expert d'IDC prévoit que la situation va déboucher, dans les années à venir, sur des amendes et des règlements à l'amiable importants. L'Union européenne sera active pour infliger des amendes aux entreprises qui enfreignent les lois sur la protection de la vie privée et autres réglementations, mais certains États américains pourraient également adopter des législations spécifiques sur l'IA, prédit-il. Les actions entreprises par les agents pourraient également enfreindre des lois américaines telles que la loi HIPAA (Health Insurance Portability and Accountability Act), qui protège la confidentialité des données médicales.

Renforcer les bonnes pratiques

Plusieurs experts en informatique et en droit estiment les prévisions d'IDC réalistes. Poursuites et amendes paraissent probables, et les plaignants n'auront pas besoin de nouvelles lois sur l'IA pour poursuivre des entreprises maîtrisant mal la technologie, estime Robert Feldman, directeur juridique du fournisseur de services de bases de données EnterpriseDB. « Si un agent d'IA cause une perte financière ou un préjudice au consommateur, les fondements juridiques existants s'appliquent déjà, note-t-il. Les autorités de réglementation peuvent intervenir dès que l'IA dépasse les limites de toute forme de conformité et de sécurité. »

Les organisations qui manquent de gouvernance et de mécanismes de contrôle rigoureux auront du mal à s'adapter à un environnement d'IA agentique, prévient Robert Feldman. « Le monde de l'IA et des données repose toujours sur les mêmes principes fondamentaux qui auraient dû régir les entreprises bien avant l'arrivée des systèmes agentiques : la responsabilité, la modération et la clarté des responsabilités, explique-t-il. Ce



que change l'IA agentique, ce n'est pas la nécessité de ces principes, mais le coût de leur non-respect. » Le passage aux agents renforcera les bonnes pratiques déjà appliquées par certaines entreprises, et d'autres constateront clairement la nécessité d'établir des garde-fous pratiques avant de déployer l'agentique à grande échelle, poursuit-il.

Les DSI joueront un rôle crucial dans la définition de ces garde-fous, ajoute Robert Feldman. « Une fois l'affaire portée devant les tribunaux, les conseils d'administration exigent des réponses sur les événements et leurs causes, explique le directeur juridique. Une explication se résumant à "c'est le système qui a fait ça" satisfait rarement un conseil d'administration. Les DSI seront de plus en plus tenus d'expliquer la gouvernance et les garde-fous mis en place pour éviter ces conséquences indésirables. »

Les DSI directement exposés

Si Ashish Nadkarni et Robert Feldman s'attendent à des amendes et à des règlements à l'amiable, les DSI doivent également craindre pour leur emploi si les agents produisent des résultats inattendus, explique Shivanath Devinarayanan, directeur technique chez

Asymbi, une entreprise spécialisée dans la modernisation du travail, à l'ère numérique. « Les poursuites judiciaires durent des années et les amendes nécessitent que les autorités de réglementation vous prennent sur le fait, dit-il. Mais un conseil d'administration peut perdre confiance en 30 secondes. Il suffit d'une question : "Que font réellement nos agents d'IA ?" Si le DSI ne peut pas répondre, c'est fini pour lui. »

Or, de nombreuses organisations semblent déployer des agents sans bien comprendre les résultats potentiels et sans politique d'IA approuvée par le conseil d'administration, selon Shivanath Devinarayanan. « La prédiction d'IDC est probablement prudente, dit-il. Le chiffre de 20 % suppose que les organisations reconnaissent le problème lorsqu'elles en ont un. La plupart ne le feront que lorsqu'il sera trop tard. »

« *Le chiffre de 20 % suppose que les organisations reconnaissent le problème lorsqu'elles en ont un. La plupart ne le feront que lorsqu'il sera trop tard. »*

Shivanath Devinarayanan

De nombreux DSI n'ont pas un inventaire complet des agents mis en production au sein de leur organisation, ajoute-t-il. « Aucun protocole d'escalade ; aucune trace d'audit entre l'action de l'agent et son impact sur l'activité, déplore le CTO. Quand un problème survient – et il y en a toujours –, les DSI n'auront aucune explication à fournir. »

Dimitri Osler, directeur technique et cofondateur de Wildix, fournisseur de communications unifiées, rétorque que les amendes semblent une conséquence probable des agents indéliçats, car les autorités de régulation n'ont pas à prouver l'intention, mais seulement que les contrôles et la gouvernance étaient inadéquats au regard du risque. L'IA n'est pas intrinsèquement dangereuse, mais de petits problèmes peuvent se transformer en catastrophes lorsqu'une organisation exploite des dizaines d'agents interdépendants, souligne-t-il. « Les systèmes agentiques peuvent agir à grande échelle et rapidement, et de petites failles de contrôle deviennent des incidents majeurs, explique-

t-il. Le problème de fond est que de nombreuses organisations adoptent l'IA parce qu'elle impressionne, et non parce qu'elle est bien gouvernée. Or, cette mentalité qui privilégie l'effet spectaculaire à la responsabilité est précisément ce qui engendre des perturbations évitables. »

Être en position de se défendre

Dimitri Osler recommande aux DSI d'adopter une approche proactive en matière de gouvernance des agents. D'exiger des preuves pour les actions sensibles et d'assurer la traçabilité de chaque action. Ils peuvent également impliquer des humains dans les tâches sensibles des agents, concevoir des agents capables de déléguer les tâches lorsque la situation est ambiguë ou risquée, et complexifier les actions à fort enjeu des agents afin de rendre plus difficile le déclenchement de mesures irréversibles, explique-t-il.

Les responsables IT devraient également investir dans la formation de leurs équipes pour mieux superviser et optimiser les agents, ajoute-t-il. « Les DSI peuvent mieux protéger les emplois en considérant la gouvernance des agents comme une pratique de leadership, et non comme une simple considération technique, affirme le cofondateur de Wildix. Organisez des exercices de simulation, définissez des règles non négociables et communiquez davantage, plutôt que de repousser le problème. »

Les DSI devraient également rendre leurs processus de gouvernance aussi transparents que possible, ajoute-t-il. « En cas d'incident, le DSI capable de présenter des contrôles clairs, des journaux d'audit, des décisions d'activation et des points de contrôle humains est bien plus en position de se défendre que celui qui se contente simplement d'affirmer que le modèle employé est responsable », conclut Dimitri Osler.

UN ARTICLE RÉDIGÉ PAR

Grant Gross, CIO US (adapté par Reynald Fléchaux)



/ TECHNOLOGIES

Évaluation et tests : le coût caché du déploiement des agents IA

Les organisations qui adoptent des agents sous-estiment souvent le coût des tests d'une technologie dont la nature non déterministe engendre fréquemment des évaluations complexes et onéreuses.

© istock

Mauvaise surprise en vue pour les entreprises ayant déployé ou déployant des agents IA ? Selon certaines études, près de 80 % des entreprises ont déjà mis en oeuvre cette technologie, mais la plupart d'entre elles ignorent le coût de leur entraînement et de l'évaluation de leurs résultats. Or, selon les experts, ces coûts peuvent largement dépasser les prévisions.

De nombreuses organisations expérimentent encore pour trouver les meilleures façons de détecter les problèmes des agents avant qu'ils ne provoquent le chaos en production, explique Lior Gavish, cofondateur et directeur technique de Monte Carlo, fournisseur de solutions d'observabilité de l'IA.

Comme beaucoup d'organisations utilisent un second LLM pour vérifier les résultats d'un agent basé sur un premier modèle de langage, les tests peuvent s'avérer bien plus coûteux que traditionnellement, ajoute-t-il. De plus, cette méthode, appelée « LLM as-a-judge », peut s'avérer plus coûteuse que le fonctionnement même de l'agent, car le coût d'utilisation d'un LLM sur une période prolongée peut rapidement devenir important.

« Tester ou de contrôler ces résultats reste complexe, souligne Lior Gavish. En pratique, on demande à un second LLM d'évaluer les performances d'un premier LLM selon divers critères, qui varient considérablement d'un cas d'utilisation à l'autre. » Monte Carlo a

elle-même constaté ce problème lorsqu'une évaluation basée sur un LLM a été menée pendant plusieurs jours, générant une facture à cinq chiffres, se remémore le directeur technique. « Une évaluation par un LLM coûte généralement beaucoup plus cher que n'importe quelle opération logicielle traditionnelle », avertit-il.

Quand les LLM évaluent les LLM

Le recours à un second LLM pour examiner les résultats d'un agent peut aussi poser un problème de confiance, car la démarche suppose que les conclusions de ce second LLM soient exactes, explique Lior Gavish. Les doutes quant à la précision des résultats peuvent engendrer de nouveaux coûts supplémentaires, les organisations menant alors d'autres tests pour les vérifier. « Ces contrôles sont non déterministes et même impossibles à reproduire, explique le directeur technique. On peut obtenir des résultats différents si l'on n'est pas vigilant. » L'approche s'éloigne donc des tests ou de la surveillance logicielle plus traditionnels, où le résultat est soit positif, soit négatif.



Il faut prendre en compte tous les tests, la journalisation et les vérifications humaines. Chaque modification implique de relancer les évaluations, et le coût grimpe très vite. »

Russell Twilligear

Le coût des évaluations des agents peut encore varier considérablement selon leur complexité, souligne Russell Twilligear, responsable de la R&D en IA chez BlogBuster, un fournisseur de contenu généré par IA. Par exemple, l'évaluation d'un agent simple et bien défini peut se limiter à quelques milliers de dollars, tandis que celle d'agents plus complexes va se chiffrer en dizaines de milliers de dollars, précise-t-il. « Il faut prendre en compte tous les tests, la journalisation et les vérifications humaines, dit-il. Chaque modification implique de relancer les évaluations, et le coût grimpe très vite. »

Indispensable contrôle humain

Les évaluations d'agents peuvent être complexes car elles doivent couvrir plusieurs facteurs, notamment le raisonnement, l'exécution, les fuites de données, le ton des réponses, la confidentialité et même l'éthique, soulignent les experts en IA. Selon Paul Ferguson, fondateur de la société de conseil Clearlead AI Consulting, une bonne évaluation intègre toujours une dimension humaine, nécessitant l'intervention d'experts du domaine pour vérifier les résultats. Il ajoute que l'un des principaux défis de l'évaluation consiste à définir ce que signifie un résultat « correct » dans des cas d'utilisation ambigus.

Dans ce type de projets, la plupart des responsables IT budgétisent les coûts évidents - temps de calcul, appels API et heures d'ingénierie -, mais négligent le coût du jugement humain nécessaire pour définir ce que Paul Ferguson appelle la « vérité de référence ». « Pour évaluer si un agent a correctement traité une requête client ou rédigé une réponse appropriée, il est indispensable que des experts du domaine évaluent manuellement les résultats et parviennent à un consensus sur ce qui constitue une réponse "correcte", note-t-il. Ce niveau d'étalonnage humain est coûteux et il est trop souvent négligé. »

L'évaluation automatisée peut être simple lorsqu'il s'agit de vérifier qu'un code franchit l'étape de compilation ou réussit tous les tests unitaires. « Mais pour les requêtes vagues comme "Aidez-moi à comprendre ces données" ou "Rédigez une réponse à ce client", définir ce qui constitue une réponse correcte devient vraiment difficile, insiste Paul Ferguson. Même les humains peuvent être en désaccord entre eux dans certains cas. »

Tester les agents : un travail de Sisyphe

Le coût élevé des évaluations provient rarement des coûts de calcul de l'agent lui-même, mais du « multiplicateur non déterministe » des tests, reconnaît Chengyu « Cay » Zhang, ingénieur logiciel et fondateur de Redcar.ai, fournisseur de solutions d'IA vocale. Il compare la formation des agents à celle des nouveaux employés, les deux étant sujettes aux aléas. « On ne

peut pas se contenter de tester un prompt une seule fois ; il faut le tester 50 fois dans différents scénarios pour vérifier la fiabilité de l'agent et éviter les erreurs, explique-t-il. Et chaque modification d'un prompt ou changement de modèle implique de relancer des milliers de simulations. »

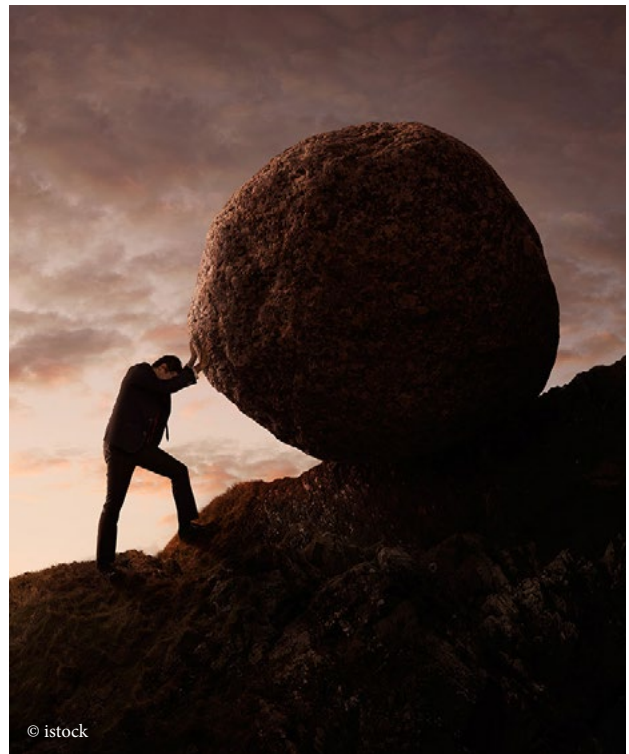
Il existe plusieurs méthodes d'évaluation des agents, notamment les tests unitaires à faible coût, l'évaluation synthétique à l'aide d'un autre modèle d'IA, les simulations d'attaques et l'observation humaine, plus onéreuse, où un expert accompagne un agent pendant une semaine ou plus, précise l'ingénieur.

Mais les organisations recherchent souvent des solutions de facilité, généralement en s'appuyant entièrement sur d'autres modèles d'IA pour l'évaluation. Un raccourci que Cay Zhang déconseille. « Pour moi, les évaluations sont une assurance, détaille-t-il. Les raccourcis dans les évaluations ne sont qu'une forme de dette technique qui se paie avec intérêts lorsque l'agent commet une erreur grossière devant un client VIP. Vous pourriez économiser 10 000 dollars sur les évaluations aujourd'hui, mais si votre agent financier commet une erreur grossière lors d'une transaction, ce coût est négligeable comparé aux dommages causés à votre image de marque. »

Limiter le champ d'action de l'agent

Si une organisation souhaite faire des économies, la meilleure solution consiste à restreindre le champ d'action de l'agent plutôt que de réduire les tests, ajoute-t-il. « Si vous négligez les étapes coûteuses, comme la vérification humaine ou les tests d'intrusion, vous vous en remettez entièrement au hasard », résume Cay Zhang.

Pour limiter les coûts d'évaluation, Paul Ferguson de Clearlead AI Consulting recommande aux organisations de commencer par des cas d'utilisation associés à des réponses claires et binaires, comme la compilation de code, avant d'aborder des scénarios plus subjectifs. Il conseille également aux organisations d'utiliser des frameworks d'évaluation de LLM, tels que LangSmith, PromptLayer ou Ragas, plutôt que de développer leur propre outillage. Selon lui, les équipes IT devraient encore commencer les tests au plus tôt.



Créer des environnements d'évaluation avant la production est bien moins coûteux que d'adapter les agents ultérieurement. »

Paul Ferguson

« Créer des environnements d'évaluation avant la production est bien moins coûteux que d'adapter les agents ultérieurement », affirme Paul Ferguson.

Lior Gavish, de Monte Carlo, propose d'autres solutions pour réduire les coûts, comme la définition de plafonds de dépenses pour les évaluations et une vérification rigoureuse des LLM utilisés pour tester les agents. « Il est possible d'optimiser légèrement le modèle, explique-t-il. Bien sûr, vous pouvez utiliser la dernière version de ChatGPT pour chaque évaluation, mais ce n'est probablement pas la solution optimale. »

UN ARTICLE RÉDIGÉ PAR

Grant Gross, CIO US (adapté par Reynald Fléchaux)