

IA INFORMATEURS ARTIFICIELS ?



ETUDE SUR LA PERCEPTION DES RÉSULTATS FINANCIERS DU CAC 40 PAR LES MODÈLES IA

Mars 2026

REPUTATION AGE

CONTEXTE

42 % des Français utilisent l'IA dans un cadre privé ★

30 % des Français utilisent l'IA dans un cadre professionnel ★

Depuis trois ans, la première page de Google n'est plus l'unique horizon. Désormais, l'enjeu est d'apparaître dans la réponse générée directement par les IA. Être cité dans ces résumés équivaut à être présenté comme une référence, parfois sans que l'utilisateur ait besoin de cliquer sur un lien. Cela transforme la gestion de la réputation...

INTERROGE LA CAPACITÉ DES LLMS À SE COMPORTER COMME DES INFORMATEURS FINANCIERS

Chaque début d'année, la saison des résultats annuels est un moment clé pour les grandes entreprises cotées. Ces publications sont scrutées par les investisseurs, analystes et journalistes.

Dans ce contexte, Reputation Age s'est mis à la place d'un investisseur individuel et a posé une question simple à ChatGPT, Claude et Gemini : quels sont les résultats financiers des sociétés du CAC 40 ?



PÉRIMÈTRE ET MODÈLES TESTÉS

39 entreprises du CAC 40

(Pernod Ricard exclue car exercice décalé)

4 indicateurs financiers demandés :

chiffre d'affaires, résultat net, résultat d'exploitation, dette nette

Modèles testés :

- ChatGPT (OpenAI - version gratuite)
- Gemini (Google - version gratuite)

ChatGPT et Gemini représentent, à une écrasante majorité (>90 %), les modèles utilisés par le grand public.

- Claude (Anthropic - version payante) ★

Claude est majoritairement utilisé dans un cadre professionnel. Claude Code est le modèle de référence pour la création d'agents

★ Seul Claude a été évalué dans sa version payante afin de se mettre dans la situation d'un professionnel avec un abonnement payant.



TROIS CHEMINS VERS L'IA

EXERCICE 1

Interrogation directe

Chaque modèle est interrogé avec un prompt identique demandant de créer un tableau contenant les quatre indicateurs financiers pour chacune des 39 sociétés;

EXERCICE 2

Interrogation par agent

L'agent RAGE a interrogé chaque modèle (156 questions par modèle) puis rempli le tableau recherché.

EXERCICE 3

Agent automatisé

L'agent RAGE est allé chercher l'information à la source (communiqués des entreprises) puis a créé le tableau recherché.

RÉFÉRENTIEL & DÉFINITION DE L'ERREUR

- Référentiel : l'agent IA développé via Claude Code a utilisé Google Search pour vérifier les données. Nous avons effectué une vérification manuelle. 20 % des résultats pourtant vérifiés via Google Search comportaient des inexactitudes et ont été corrigés avant utilisation. C'est notre base de comparaison.
- Seuil d'erreur : il y a erreur lorsqu'il y a écart $\geq 5\%$ par rapport à la donnée de référence.
- Données manquantes : Absence de réponse traitée comme une erreur à part entière.
- Confusion : une confusion avec un exercice précédent (ce point n'a pas été vérifié systématiquement) est également traitée comme une erreur à part entière.



EXERCICE 1

Interrogation directe : une tâche à la limite des capacités des chatbots

CHATGPT

OpenAI

0 %

de réponses exactes

ChatGPT n'a pas pu générer de tableau

GEMINI

Google

12,9 %

de réponses exactes

20 correctes

136 erronées

sur 156 réponses analysées

CLAUDE

Anthropic ★

7 %

de réponses exactes

11 correctes

145 erronées

sur 156 réponses analysées

★ Seul Claude a été évalué dans sa version payante afin de se mettre dans la situation d'un professionnel avec un abonnement payant.



EXERCICE 2

Interrogation par agent : aucun modèle ne dépasse 15% de réponses exactes

CHATGPT

OpenAI

9,6%

de réponses exactes

15 correctes

141 erronées

sur 156 réponses analysées

GEMINI

Google

12,8 %

de réponses exactes

20 correctes

136 erronées

sur 156 réponses analysées

CLAUDE

Anthropic

7,1 %

de réponses exactes

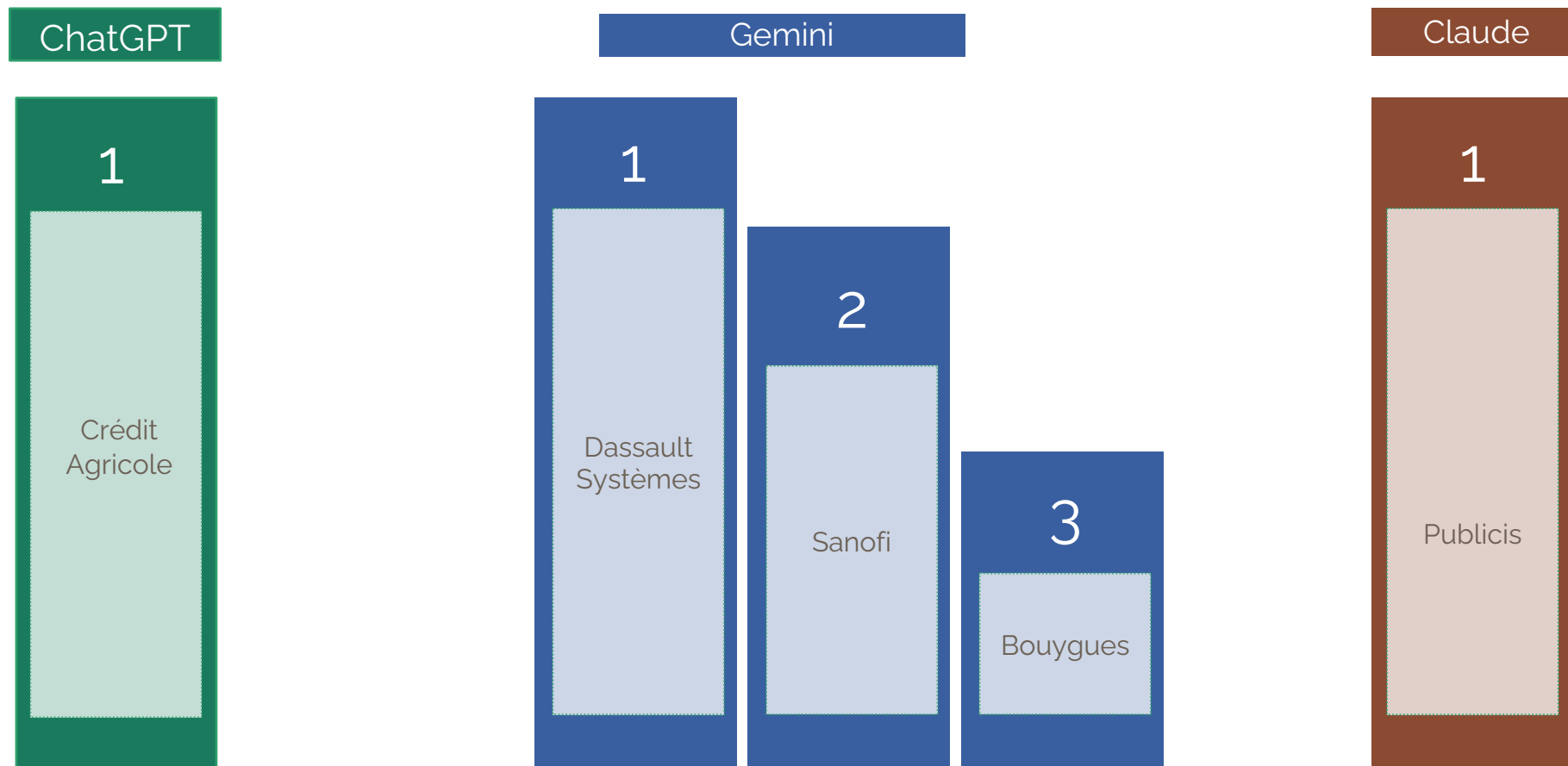
11 correctes

145 erronées

sur 156 réponses analysées

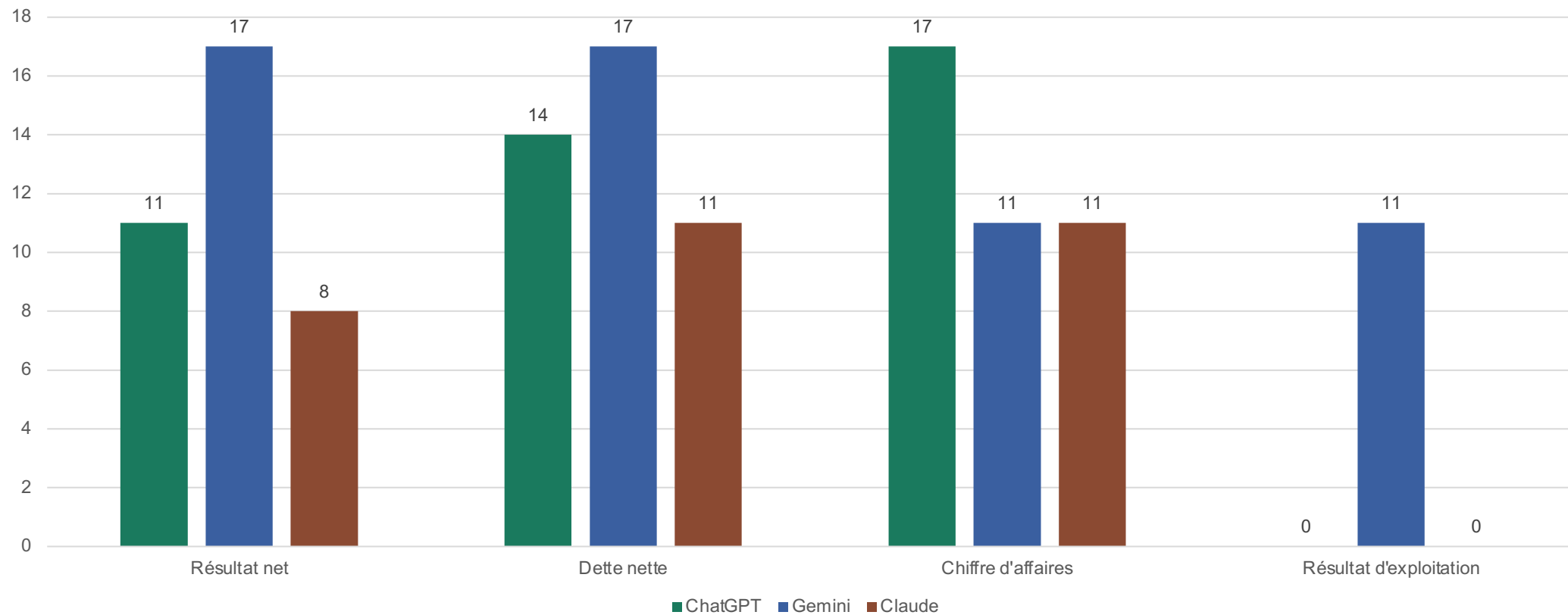


GEO : les entreprises les mieux représentées



RÉSULTATS PAR INDICATEUR

Taux de réponses exactes (%) selon l'indicateur demandé



Agent automatisé :
Encore un effort

C L A U D E

Anthropic

80 %

de réponses exactes

125 correctes

31 erronées

sur 156 réponses analysées



Ce que révèle l'étude

EXERCICE 1 — INTERROGATION DIRECTE

- ChatGPT génère un tableau vide
- Gemini répond rapidement avec des erreurs
- Claude ne répond que dans la version payante

EXERCICE 2 — INTERROGATION PAR AGENT

- Résultats très décevants pour la quasi-totalité des modèles
- Les données restituées présentent des écarts notables avec les chiffres officiels
- De larges disparités entre modèles.
- On peut constater que la communication officielle des sociétés n'est pas une référence pour les LLMs.

EXERCICE 3 — AGENT AUTOMATISÉ

- Il y a une prime aux payeurs pour obtenir des informations fiables.
- Une importante quantité de données n'est pas pour autant synonyme de précision.
- 20 % des données ont nécessité une correction manuelle. Ces erreurs peuvent être attribuées à des défauts d'architecture dans les supports de communication des sociétés.
- Les entreprises doivent s'adapter à ces nouveaux informateurs.

Ce qu'il faut retenir

01 Les approximations dangereuses

Pour les modèles grand public, plus de 8 réponses sur 10 sont inexactes. La précision financière reste hors de portée pour l'IA.

02 Performances quasi identiques entre modèles

ChatGPT, Claude et Gemini se situent dans un même palier de faiblesse.

03 La communication officielle des sociétés n'est pas une référence pour les LLMs

Les LLM ne prennent pas les sites officiels comme unique référence.



C O N C L U S I O N

Et demain pour les entreprises ?

Les IA ne vont pas chercher les informations auprès des sources propres.

Là où les entreprises du CAC 40 innovent et trouvent de nouveaux marchés, ces mêmes entreprises semblent en retard dans leur communication financière à l'heure des IA.

Comment vont-elles rattraper ce retard ?



Nos pistes de solutions : notre savoir-faire GRO

Nous travaillons votre GRO (Generative Reputation Optimization).

Nous analysons :

- La « part de modèle » de votre entreprise dans les IA
- Les messages dominants associés à votre marque (innovation, ESG, risques, controverses...)
- Les écarts entre votre discours corporate et la perception générée par l'IA
- Votre positionnement face à vos concurrents dans les réponses des IA

Nous produisons :

- Un audit IA de votre réputation avec un benchmark avec vos concurrents.
- Des recommandations opérationnelles : contenus et sources à renforcer et retravailler, stratégie GRO pour influencer les réponses IA, amélioration de la structure du site.



Et demain, notre Observatoire IA de la communication financière

- Une analyse des performances GEO des grands groupes avec des pistes de recommandations
- ChatGPT / Claude / Gemini – analyse quadrimestrielle en tenant compte des calendriers financiers
- Un scoring par entreprise
- Une dimension qualitative au-delà des chiffres
- Une analyse des performances GEO des grands groupes avec des pistes de recommandations



Contact

Charles-Henri D'Auvigny
Managing Partner
charles-henri@reputation-age.com
06 09 67 49 81

Hamza Benmira
Data Analyst
hamza@reputation-age.com
06 05 54 40 56

Karl Gateau
Consultant IA & Réputation
karl@reputation-age.com
06 51 95 29 39





REPUTATION AGE

L'étude a été réalisée par Charles-Henri d'Auvigny, Hamza Benmira et
Karl Gateau.

