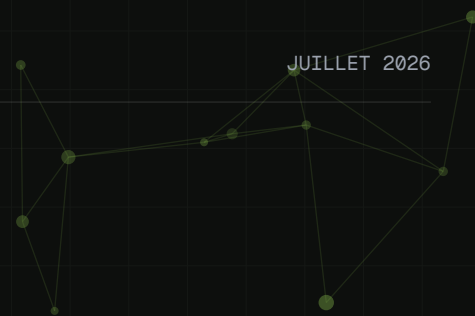


ALGORITHME LINKEDIN 2026



Ce qui est prouvé, ce qui est mesuré,
ce qui est inventé

Un état des lieux des sources publiques. Sans conseils, sans tactiques.

Prouvé

LinkedIn l'a écrit ou publié : article de recherche, blog d'ingénierie, communication officielle.

Mesuré

Une étude indépendante, avec méthodologie et volume de données publiés.

Contesté

Plusieurs sources sérieuses arrivent à des conclusions incompatibles.

Inventé

Circule abondamment, ne remonte à rien. La précision du chiffre est le seul argument.

*Ce document ne recommande rien. Il établit ce qui est connu,
et comment on peut le savoir.*

NAVIGATION

SOMMAIRE

Ce document compare trois choses : ce que LinkedIn publie, ce que les études indépendantes mesurent, et ce que tout le monde répète.

01	En résumé	03
02	Pourquoi un état des lieux plutôt qu'un guide ?	04
03	Que publie réellement LinkedIn ?	05
04	Le fil est-il classé par une intelligence artificielle ?	06
05	Comment le système comprend-il ce dont parle une publication ?	08
06	Que cherche réellement le système de classement ?	10
07	Les liens externes sont-ils pénalisés ?	11
08	Que valent les chiffres qui circulent ?	13
09	Ce que personne ne sait	14
10	Comment évaluer une source sur ce sujet ?	15
11	Note de méthode et limites	16
12	Questions fréquentes	17
13	Glossaire	18
14	Sources	19

Un état des lieux des sources publiques. Sans conseils, sans tactiques.

SECTION 01

EN RÉSUMÉ

Le 12 février 2026, vingt-quatre ingénieurs de LinkedIn ont publié un article décrivant le modèle qui classe le fil de 1,2 milliard de membres. Il est en accès libre. Presque aucun des articles qui expliquent « le nouvel algorithme » ne le cite.

Ce document compare trois choses : ce que LinkedIn publie, ce que les études indépendantes mesurent, et ce que tout le monde répète. Les trois ne se recoupent pas.

Quatre constats en ressortent.

Le plus dérangeant d'abord, parce qu'il contredit l'affirmation la plus répandue de l'année : dans les deux documents que LinkedIn a publiés au premier trimestre 2026, le classement du fil n'est pas assuré par une intelligence artificielle générative. L'entreprise en a construit une et l'a testée en conditions réelles. À la date de son article, elle n'obtenait pas de meilleurs résultats que le modèle déjà en service, et n'était pas en production.

Deuxième constat, plus simple : le système ne cherche à provoquer qu'un petit nombre de comportements précis, que LinkedIn nomme un par un dans ses publications.

Le troisième oblige à l'inconfort. Sur les liens externes, trois études sérieuses aboutissent à des conclusions incompatibles, dont une inverse aux deux autres, si bien qu'aucune réponse ferme n'est possible en l'état. Reste le quatrième, plus simple à énoncer : plusieurs chiffres très diffusés ne remontent à rien.

Ce document ne recommande rien. Il établit ce qui est connu, et comment on peut le savoir.

SECTION 02

POURQUOI UN ÉTAT DES LIEUX PLUTÔT QU'UN GUIDE ?

Il existe déjà des centaines de guides sur l'algorithme LinkedIn. Il n'existe presque rien sur la fiabilité de ce qu'ils affirment.

Le contraste est saisissant. D'un côté, LinkedIn publie ses travaux dans des revues scientifiques, avec le détail de ses modèles et les résultats de ses tests. C'est gratuit et ouvert à tous. De l'autre, ce qu'on lit sur le sujet vient de billets qui commentent d'autres billets, chacun ajoutant une couche de certitude à une information dont plus personne ne connaît l'origine.

Résultat : le vérifiable et l'inventé se lisent exactement pareil. Nous avons voulu séparer les deux. La méthode est simple : remonter à la source primaire de chaque affirmation, ou constater qu'elle n'en a pas.

— Les quatre niveaux de preuve utilisés dans ce document

01

Prouvé

LinkedIn l'a écrit ou publié : article de recherche, blog d'ingénierie, communication officielle. Niveau de confiance le plus élevé, avec une réserve : une entreprise décrit ce qu'elle veut bien décrire.

02

Mesuré

Une étude indépendante avec méthodologie et volume de données publiés. Fiable, mais l'échantillon n'est jamais représentatif de l'ensemble de la plateforme.

03

Contesté

Plusieurs sources sérieuses arrivent à des conclusions incompatibles. Il faut le dire, pas trancher.

04

Inventé

Circule abondamment, ne remonte à rien. La précision du chiffre est souvent le seul argument.

SECTION 03

QUE PUBLIE RÉELLEMENT LINKEDIN ?

Tout repose sur trois publications. Deux d'entre elles ne sont presque jamais citées dans ce qu'on lit sur LinkedIn.

La plus importante est un article de recherche de février 2026, *An Industrial-Scale Sequential Recommender for LinkedIn Feed Ranking*, déposé sur arXiv. Vingt-quatre ingénieurs le signent. Son résumé indique que ce modèle est alors le système principal du fil. L'article décrit ensuite ses objectifs, son architecture et les résultats de ses essais, avec un niveau de détail qu'aucune autre source n'approche. C'est le document dont on ne trouve pratiquement aucune trace dans ce qui s'écrit sur le sujet.

Vient ensuite un article de mars 2026 publié sur le blog technique de LinkedIn, signé Hristo Danchev, consacré cette fois à la présélection : comment la plateforme choisit les quelques centaines de publications qu'elle va ensuite ordonner. C'est celui que la presse spécialisée a repris, souvent sans le lire.

Restent deux documents plus anciens. Une note technique de 2020 sur le temps de lecture, régulièrement citée de travers. Et *LiRank*, présenté dans une conférence scientifique en 2024, qui décrit le système précédent.

Ces quatre documents ont été lus intégralement. C'est la seule façon de constater à quel point les reprises qui en circulent sont partielles.

SECTION 04

LE FIL EST-IL CLASSÉ PAR UNE INTELLIGENCE ARTIFICIELLE ?

C'est l'affirmation la plus répandue de 2026. Les deux publications de LinkedIn la contredisent.

— Ce qui est prouvé

Dans les deux documents publiés au premier trimestre 2026, le classement du fil est assuré par un modèle que le blog d'ingénierie appelle Generative Recommender et que l'article de recherche nomme Feed SR. C'est un transformeur séquentiel à attention causale, autrement dit un modèle qui lit vos interactions dans leur ordre chronologique sans jamais anticiper la suite. Il a remplacé un modèle d'une génération antérieure, d'une famille technique différente. Ce n'est pas un modèle de langage.

Les modèles de langage interviennent bien, mais à deux endroits qui ne sont pas le classement : la récupération, c'est-à-dire la présélection des publications candidates, et la représentation du profil du membre qui consulte son fil. En clair : un modèle de langage décide quelles publications sont candidates. Un modèle beaucoup plus petit décide dans quel ordre elles apparaissent.

— D'où vient le malentendu

D'une phrase de LinkedIn elle-même, ce qui rend l'erreur assez compréhensible. Le deuxième paragraphe du blog de mars 2026 annonce le déploiement, dans les termes exacts de LinkedIn, d'« a new advanced ranking system, powered by LLMs and GPUs ». C'est la phrase que tout le monde a citée, et elle n'est pas fausse : le système comporte bien un modèle de langage.

Le corps du même article dit pourtant autre chose, en détail et sur plusieurs pages. La présélection repose effectivement sur un modèle de langage réentraîné pour cet usage, monté en double encodeur, c'est-à-dire une méthode qui traduit séparément le membre et la publication en deux descriptions comparables. Le classement, lui, repose sur un tout autre modèle, un transformeur séquentiel qui confie chaque type d'action à un sous-modèle spécialisé. Les deux étapes sont bien distinctes, et l'article les sépare. Ce qui a fait le tour du web, c'est la phrase d'introduction.

— Ce que LinkedIn dit avoir essayé

L'article consacre une section entière aux architectures alternatives évaluées avant le choix final. L'équipe a réellement construit un classeur fondé sur un modèle de langage. Chaque publication était décrite en texte dans une requête, et le modèle répondait à des questions du type « ce membre va-t-il cliquer ? ». L'article expose trois obstacles rencontrés.

La première tient à la nature des données. Un modèle de langage travaille sur du texte, or une bonne part de ce qui compte pour classer une publication ne s'exprime pas en mots : le nombre de réactions, l'ancienneté, la proximité avec l'auteur. Traduites en phrases, ces valeurs perdaient leur utilité, et le modèle ne parvenait pas à les exploiter correctement.

Vient ensuite le coût. Un modèle de langage découpe le texte en petits morceaux qu'il facture à l'unité. Chaque publication en consommait des centaines, contre deux dans Feed SR.

Et puis il y a l'obstacle décisif, celui de la performance. Au moment où l'article est écrit, ce classeur n'obtenait pas de meilleurs résultats en ligne que le modèle déjà en service. Le résumé de l'article l'énonce sans détour : les auteurs expliquent « why Feed-SR provided the best combination of online metrics and production efficiency ». Il se débrouillait honorablement sur les contenus hors réseau, et échouait sur ceux du réseau, la force d'une relation se laissant mal décrire en mots.

Une précision, ici, et elle vaut pour tout ce document. LinkedIn décrit un résultat obtenu à une date, pas une décision définitive. L'entreprise n'a annoncé aucun abandon, et rien n'indique ce qu'il en est depuis. Ce qui est établi tient en une phrase : le système décrit au premier trimestre 2026 classe le fil avec un transformeur séquentiel, et la voie du modèle de langage n'avait alors pas donné de meilleurs résultats.

— D'où vient alors l'histoire du modèle géant ?

De 360Brew, et l'histoire vaut d'être racontée parce qu'elle montre comment une information se déforme sans que personne ne mente vraiment.

360Brew existe bel et bien. C'est un article de recherche déposé sur arXiv en janvier 2025 par une équipe de LinkedIn, qui décrit un modèle de 150 milliards de paramètres capable en principe de traiter d'un seul tenant le classement du fil, les recommandations d'emploi et les suggestions de relations. Son résumé le qualifie en toutes lettres de « research pre-production model », autrement dit de prototype non déployé.

Un an plus tard, en janvier 2026, Forbes publie un entretien dans lequel un créateur suivi par 1,2 million de personnes affirme que la dernière mise à jour de LinkedIn s'appelle 360 Brew et qu'elle « now shows your content more accurately to your ICP ». En quelques semaines, l'affirmation devient un fait acquis, et des dizaines de publications construisent sur cette base des recommandations tactiques détaillées.

Trois éléments la contredisent pourtant, tous vérifiables en quelques minutes. L'article se présente lui-même comme pré-production. Nous n'avons trouvé aucune confirmation de déploiement par LinkedIn. Et sa publication technique de mars 2026 ne mentionne nulle part ce nom.

Un quatrième élément permet de mesurer combien la chaîne de recopiage s'est éloignée de la source. Plusieurs articles affirment que 360Brew repose sur le modèle LLaMA 3, quand l'article de recherche décrit une architecture Mixtral. L'erreur porte sur une information qui figure en clair dans le texte : elle indique que la source n'a pas été consultée avant d'être citée. En douze mois, un modèle de laboratoire est ainsi devenu le nom couramment donné à l'algorithme de LinkedIn.

SECTION 05

COMMENT LE SYSTÈME COMPREND-IL CE DONT PARLE UNE PUBLICATION ?

Le blog d'ingénierie décrit ici un mécanisme que personne ne reprend, et qui éclaire pourtant tout le reste. Sous l'intitulé « Unified retrieval through fine-tuned LLMs », le blog explique que LinkedIn convertit chaque publication et chaque membre en un texte, envoyé à un modèle de langage qui en produit une représentation mathématique. La comparaison ne se fait plus sur des mots-clés, mais sur du sens.

Ce que contient le texte décrivant une publication est explicitement listé : son format, son texte, ses compteurs de réactions et de vues, les informations de l'article s'il y en a un, et celles de son auteur, à savoir son nom, son titre de profil, son entreprise et son secteur. Autrement dit, le titre de profil de l'auteur fait partie de la description de sa publication.

Côté membre, le texte rassemble le profil, les compétences, le parcours, la formation, et la liste chronologique des publications avec lesquelles la personne a interagi.

— L'erreur que LinkedIn raconte lui-même

LinkedIn raconte une erreur, et le récit vaut mieux qu'un long développement sur les limites des modèles de langage. La section « From structured data to effective prompts » raconte l'épisode. Les compteurs d'engagement étaient d'abord transmis bruts. Une publication à 12 345 vues apparaissait dans le texte sous la forme « views:12345 ». Le modèle traitait ces chiffres comme n'importe quels caractères, et le résultat a inquiété les équipes : la corrélation entre la popularité d'une publication et sa pertinence calculée était de moins 0,004, autrement dit nulle.

Or la popularité est un signal fort. Une publication que beaucoup de gens trouvent utile a de bonnes chances de l'être aussi pour vous.

La correction a consisté à remplacer le nombre brut par sa position dans le classement. « views:12345 » est devenu « 71e centile », c'est-à-dire : cette publication est plus vue que 71 % des autres. Le modèle a enfin pu apprendre ce que veut dire « au-dessus de la moyenne ».

La corrélation a été multipliée par trente. La qualité de la présélection a progressé de 15 %.

— Ce que le système apprend de vos silences

Deux détails d'entraînement méritent d'être connus, parce qu'ils décrivent ce que la machine considère comme un signal. Le premier, exposé sous le titre « Training dual encoders at scale », concerne les contre-exemples. Pour apprendre à distinguer une publication vaguement pertinente d'une publication réellement utile, le modèle est entraîné sur des publications qui ont été affichées à un membre sans provoquer la moindre réaction. LinkedIn les appelle des négatifs difficiles. En ajouter deux par membre a amélioré la qualité de la présélection de 3,6 %.

Le second est plus contre-intuitif. L'historique d'un membre incluait au départ tout ce qui lui avait été montré, y compris ce qu'il avait fait défiler sans s'arrêter. Les équipes ont retiré ce dernier ensemble. Le modèle a mieux appris, et l'entraînement est devenu 2,6 fois plus rapide.

La formulation de LinkedIn mérite d'être citée :

le modèle apprend mieux de ce que vous avez choisi que de tout ce qu'on vous a montré.

— Le cas du membre qui vient d'arriver

Le blog décrit un scénario précis, qu'il désigne par le terme « cold-start » : un nouveau membre inscrit, dont on ne connaît que le titre de profil et l'intitulé de poste, sans aucun historique. Le modèle de langage déduit de ces deux seules informations les sujets susceptibles de l'intéresser, grâce à ce qu'il a appris sur le monde en lisant des milliards de textes. LinkedIn donne un exemple : quelqu'un dont le profil indique « ingénierie électrique » mais qui interagit avec des contenus sur les petits réacteurs modulaires. Un système par mots-clés ne fait pas le lien. Un modèle de langage sait que ces sujets se touchent.

SECTION 06

QUE CHERCHE RÉELLEMENT LE SYSTÈME DE CLASSEMENT ?

L'article de février 2026 répond sans détour, et c'est sans doute l'information la plus utile de tout ce dossier. Le modèle prédit six actions, et le blog d'ingénierie les nomme toutes.

La section « Teaching transformers to recommend » les énumère. Trois sont passives : le clic, le passage sans arrêt, et le **Long Dwell**, c'est-à-dire le fait de rester sur une publication au-delà d'un seuil qui dépend de son type. Trois sont actives : la réaction, le commentaire et le partage. L'article de recherche regroupe ces trois dernières sous le nom de **Contribution**.

Ces deux familles ne sont pas mélangées. Le modèle les traite séparément au moment de prédire, en s'appuyant sur la même lecture de votre historique.

Les deux mesures que LinkedIn met en avant dans son article de recherche, celles sur lesquelles il rapporte ses gains, sont le Long Dwell et la Contribution.

— Ce que l'article documente aussi, chiffres à l'appui

Sa mémoire est bornée : mille publications vues, prélevées sur environ un an. Au-delà, plus rien. Le poids de chacune décroît avec le temps, selon deux mécanismes distincts. Les données servant à entraîner le modèle perdent la moitié de leur influence tous les soixante jours. Et dans la séquence d'un même membre, les interactions les plus anciennes comptent environ moitié moins que les plus récentes. Un historique se périmé, à un rythme mesurable.

Autre signal, moins attendu : la façon dont les lecteurs précédents se répartissent par tranche de temps de lecture, de zéro à cinq secondes jusqu'à plus de soixante. Ce seul élément améliore de 2,5 points la prédiction du Long Dwell.

LinkedIn corrige aussi un biais que peu de gens soupçonnent. Une publication placée en tête du fil récolte mécaniquement plus d'interactions qu'une autre, quelle que soit sa qualité, simplement parce qu'elle est vue en premier. Le système compense en pondérant chaque observation selon la probabilité qu'elle avait d'être vue, et en apprenant un décalage propre à chacune des soixante premières positions. Autrement dit, il essaie de mesurer ce que vaut vraiment une publication, indépendamment de la chance qu'elle a eue d'être bien placée.

Résultat de l'essai grandeur nature : 2,10 % de temps passé en plus, avec les gains les plus élevés chez les membres les plus actifs et aucun effet mesurable chez les nouveaux.

— Une conséquence pour lire le reste du marché

Une affirmation courante veut que le fait d'enregistrer une publication, c'est-à-dire de la mettre de côté pour la relire plus tard, soit devenu le premier facteur de classement. L'article ne les cite pas parmi les objectifs optimisés. Cela ne prouve pas qu'ils ne comptent pas, mais l'affirmation ne s'appuie sur aucune source primaire, alors que les deux objectifs nommés, eux, sont documentés.

SECTION 07

LES LIENS EXTERNES SONT-ILS PÉNALISÉS ?

C'est la question la plus posée, et celle où l'honnêteté impose de ne pas conclure.

— La position de LinkedIn

Nous n'avons trouvé, dans la documentation publique de LinkedIn, aucune règle pénalisant les publications contenant un lien. Sa documentation décrit des critères de qualité et de sécurité, et réduit la distribution des contenus jugés de faible qualité ou relevant du spam. Les liens n'y figurent pas.

— Ce que mesurent les études

Elles ne s'accordent pas, et l'écart est considérable. Une étude publiée par Ordinal, portant sur plus de 900 000 publications réparties sur trente-sept mois, rapporte une baisse de portée de 26,5 % pour les publications contenant un lien, avec un test statistique de Mann-Whitney et une valeur p inférieure à 0,001, ce qui signifie moins d'une chance sur mille que l'écart observé soit dû au hasard. C'est la seule étude disponible à publier ce genre de vérification statistique.

Deux réserves l'accompagnent, et nous les signalons parce que notre propre grille de lecture les impose. Ordinal édite un outil de publication LinkedIn dont l'une des fonctions consiste à déplacer automatiquement les liens en premier commentaire : l'entreprise a donc un intérêt commercial à ce que la pénalité soit réelle. Et sur une autre page de son site, elle présente cette fonction comme un moyen d'éviter « the 60% external link penalty », c'est-à-dire précisément le chiffre sans source que nous rangeons plus loin parmi les inventions.

Le rapport *Algorithm Insights* de Richard van der Blom, fondé sur 1,3 million de publications, rapporte une baisse de 18,8 % de la portée médiane, c'est-à-dire celle de la publication qui se situe au milieu du classement, pour un lien placé dans le corps du texte. Le rapport *State of the Algorithm* du premier trimestre 2026, publié par Saywhat à partir de 397 605 publications, aboutit à l'inverse : les publications contenant plusieurs liens externes y performant nettement mieux que celles qui n'en contiennent aucun. Les auteurs attribuent ce résultat à la qualité du contenu, les publications riches en ressources utiles générant des signaux d'engagement suffisants pour compenser.

Quant au chiffre de 60 % de pénalité, largement diffusé, il ne remonte à aucune source identifiable.

— L'hypothèse qui réconcilie ces résultats

Elle découle de ce que l'on sait du modèle. À supposer qu'aucun détecteur de lien ne sanctionne les publications, un enchaînement mécanique suffit à produire le même résultat.

Un lecteur clique, il quitte la plateforme. Il ne lit donc plus la publication, ce qui effondre le Long Dwell. Il ne commente ni ne partage, ce qui supprime la Contribution. Le modèle observe un contenu qui ne retient pas et ne fait pas réagir, et le distribue moins.

Le lien ne serait donc pas puni : il produirait mécaniquement les signaux d'un contenu médiocre, ce qui revient au même du point de vue du modèle mais change tout dans l'interprétation. Il y a aussi une explication plus terre-à-terre, dont personne ne parle. Beaucoup de publications à lien sont creuses, publiées uniquement pour placer une adresse. Elles sous-performent parce qu'elles sont vides.

Reste que cette hypothèse n'est pas démontrée. Elle est cohérente avec l'ensemble des observations, ce qui n'est pas la même chose.

SECTION 08

QUE VALENT LES CHIFFRES QUI CIRCULENT ?

Nous avons retenu quatre affirmations parmi les plus diffusées, choisies parce qu'elles sont invérifiables et qu'elles portent toutes un chiffre précis.

— Les taux d'engagement par tranche de temps de lecture

On lit régulièrement que l'engagement passe de 1,2 % en dessous de trois secondes à 15,6 % au-delà de soixante secondes. Les tranches existent, l'article de février 2026 le confirme. Mais nous n'avons trouvé aucun taux d'engagement par tranche dans les publications de LinkedIn. Ces nombres sont même repris par des sites qui, sur une autre de leurs pages, signalent qu'ils ne sont pas traçables.

— Le « Depth Score »

Présenté depuis fin 2025 comme le nouveau critère de LinkedIn. Ce terme n'apparaît dans aucune des publications de LinkedIn que nous avons consultées, et nous n'avons trouvé aucune source qui l'attribue explicitement à l'entreprise. L'idée derrière ce mot, mesurer combien de temps on s'arrête sur un contenu, est bien réelle et documentée depuis 2020 sous le nom de temps de lecture. Le nom, lui, semble venir d'ailleurs.

— Les trois pourcentages qui circulent en français

On les retrouve dans de nombreux articles et posts francophones, toujours sans référence. Commenter sous sa propre publication avant tout autre commentaire coûterait 20 % de portée. Répondre dans la première heure en ferait gagner 35 %. Commenter ailleurs quinze à trente minutes avant de publier en ajouterait 21 %.

Nous n'avons retrouvé aucun de ces trois chiffres, ni dans les publications de LinkedIn, ni dans une étude dont la méthode et l'échantillon soient publiés.

— Les taux d'engagement par format

Ils sont irréconciliables entre eux. Pour les carrousels, selon la source, on trouve 1,44 %, 6,60 % ou 24,42 %. Aucun de ces outils ne publie sa méthode de calcul : on ignore ce qu'il place au dénominateur, et sur quel échantillon il travaille. Faute de cette information, les comparer entre eux n'a pas de sens.

— Un mot sur la façon dont ces chiffres se propagent

Un chiffre précis se partage mieux qu'une incertitude honnête, et il ne demande jamais à être vérifié. C'est probablement la raison pour laquelle « moins 18,8 % selon une étude portant sur 1,3 million de publications » circule moins bien que « moins 60 % ».

SECTION 09

CE QUE PERSONNE NE SAIT

Un état des lieux honnête doit aussi dire ce qu'on ne sait pas. Le poids exact de chaque facteur ne sera jamais public. Quand vous lisez « tel critère compte pour X pour cent », c'est une invention.

On ignore aussi ce qui s'est passé depuis mars 2026. LinkedIn a continué de publier, mais sur d'autres sujets : son assistant de recrutement en juin 2026, ses outils internes de qualité logicielle. Aucune publication consacrée au classement du fil n'est venue depuis celle de mars, du moins aucune que nous ayons trouvée. Parler de l'algorithme « actuel » suppose donc que rien d'important n'a changé en quelques mois. C'est probable. Ce n'est pas vérifiable.

Dans le détail, plusieurs points restent ouverts. Les seuils de Long Dwell varient selon le type de publication, l'article le confirme, mais aucune valeur n'est donnée, ce qui interdit toute règle de longueur fondée sur autre chose que ses propres résultats. Le sort des liens placés en commentaire reste franchement incertain, certaines données suggérant une dégradation et d'autres non. Et l'effet d'une modification après publication n'a jamais été mesuré, alors que la croyance qu'elle détruit la portée est répandue partout.

Les études indépendantes que nous avons consultées reposent toutes sur les clients de l'outil qui les publie, et aucune ne dit travailler sur un échantillon tiré au hasard parmi tous les membres. Leurs résultats sont utiles, et faussés dès le départ.

Sur un dernier point, il faut croire LinkedIn sur parole. L'entreprise dit contrôler régulièrement ses modèles pour vérifier deux choses : que les publications d'auteurs différents sont traitées de la même façon, et que ce que voient les uns ressemble à ce que voient les autres. Elle ajoute que ses modèles regardent l'activité professionnelle et jamais l'âge, le sexe ou l'origine. Ces contrôles ne sont pas publiés.

SECTION 10

COMMENT ÉVALUER UNE SOURCE SUR CE SUJET ?

Quatre questions suffisent à écarter l'essentiel du bruit, et elles ne demandent aucune compétence technique.

— La source cite-t-elle une publication primaire vérifiable ?

Un lien vers arXiv ou vers le blog d'ingénierie de LinkedIn, pas vers un autre article de blog. Une source qui ne cite jamais arXiv n'a pas lu les travaux de LinkedIn.

— Donne-t-elle son volume de données et sa période ?

Sans échantillon ni fenêtre temporelle, un pourcentage ne vaut rien.

— Vend-elle quelque chose que sa conclusion sert ?

Un éditeur d'outil de carrousels qui conclut à la supériorité des carrousels n'est pas un observateur neutre.

— Le chiffre est-il trop précis pour son origine ?

Une valeur issue d'une étude nommée et datée peut être précise. Une affirmation générale assortie d'un pourcentage rond, sans source, ne l'est jamais légitimement.

SECTION 11

NOTE DE MÉTHODE ET LIMITES

Ce document a été établi en juillet 2026 et porte sur l'état public des connaissances à cette date.

Un mot sur le vocabulaire employé. Quand nous écrivons qu'une information n'existe pas, il faut lire que nous ne l'avons pas trouvée dans les sources listées en fin de document. Prouver l'absence d'une chose est impossible. Signaler qu'on ne la trouve nulle part est vérifiable par quiconque refait le chemin.

Les publications de LinkedIn décrivent des systèmes et des résultats de tests, jamais le poids exact de chaque critère. Ce qui y figure est donc exact et forcément incomplet, et nous n'avancons rien au-delà de ce qu'elles écrivent. Ce document vieillira, aussi. Ce qu'il décrit date de février et mars 2026, la plateforme évolue vite, et les principes résisteront mieux que les chiffres.

Nous n'avons traité ni la publicité, ni les pages entreprise, ni les recommandations d'emploi. Ce sont des systèmes différents, qui méritent un examen séparé.

SECTION 12

QUESTIONS FRÉQUENTES

LinkedIn publie-t-il vraiment le fonctionnement de son algorithme ?

En partie, oui. L'entreprise publie le détail de ses modèles, ses mesures et les résultats de ses essais sur de vrais utilisateurs, dans des revues scientifiques et sur son blog technique. Elle ne publie pas les pondérations de ses modèles. Ce qui est décrit est donc réel, mais forcément incomplet.

Pourquoi si peu d'articles citent-ils ces publications ?

Elles sont techniques et longues, donc mal adaptées au format court. Il est plus rapide de commenter un billet de blog déjà commenté ailleurs. Le prix de cette facilité : au bout de quelques recopiations, plus personne ne sait d'où vient l'information.

Les études indépendantes sont-elles fiables ?

Elles sont utiles, et faussées dès le départ. Celles que nous avons consultées portent sur les clients de l'outil qui les publie, et aucune n'indique travailler sur un tirage au hasard parmi tous les membres. Quand elles comparent deux situations dans le même groupe, on peut s'y fier. Quand elles annoncent un chiffre absolu, beaucoup moins.

Que faut-il retenir si l'on ne retient qu'une chose ?

Que le système de classement prédit six actions nommées, regroupées en deux familles, et que toute affirmation qui n'agit sur aucune d'entre elles mérite d'être questionnée.

Le titre de profil de l'auteur influence-t-il la diffusion de ses publications ?

Le blog d'ingénierie liste ce qui compose la description d'une publication envoyée au modèle : son texte, son format, ses compteurs, et les informations de son auteur, dont son titre de profil, son entreprise et son secteur. Le titre fait donc partie de ce qui sert à situer un contenu. LinkedIn ne quantifie pas cet effet.

Ce document donne-t-il des recommandations ?

Non, et c'est délibéré. Il établit ce qui est connu et comment on peut le savoir. Ce qu'il faut en faire dépend de chaque situation.

SECTION 13

GLOSSAIRE

Les termes ci-dessous reviennent dans les publications de LinkedIn. Ils sont donnés dans le sens exact que l'entreprise leur donne, pas dans un sens général.

AUC. Mesure standard de la qualité d'une prédiction, comprise entre 0,5 et 1. Une AUC de 0,5 équivaut au hasard. Les gains rapportés dans ce domaine se comptent en fractions de point, ce qui explique qu'un gain de deux points soit présenté comme important.

Centile. Position d'une valeur dans une distribution. Une publication au 71e centile de vues est plus vue que 71 % des autres. LinkedIn a remplacé les compteurs bruts par des centiles pour les rendre exploitables par ses modèles.

Contribution. Terme employé par LinkedIn pour désigner l'ensemble formé par la réaction, le commentaire et le partage. C'est l'un des deux objectifs sur lesquels l'entreprise rapporte ses gains.

Double encodeur. Architecture dans laquelle un même modèle traite séparément deux objets, ici un membre et une publication, pour produire deux représentations directement comparables entre elles.

Feed SR, ou Generative Recommender. Le modèle de classement en production. Deux noms pour un même système : le premier vient de l'article de recherche, le second du blog d'ingénierie.

Long Dwell. Le fait de rester sur une publication au-delà d'un seuil de temps, seuil qui varie selon le type de publication. Second objectif sur lequel LinkedIn rapporte ses gains. À ne pas confondre avec le temps de lecture, qui est la mesure brute dont le Long Dwell est un seuil.

Modèle de langage. Modèle entraîné sur de très grands volumes de texte, capable d'établir des liens de sens qu'aucune règle explicite ne lui a appris. Chez LinkedIn, il sert à la présélection des publications, pas à leur classement.

Négatifs difficiles. Publications réellement affichées à un membre et qui n'ont provoqué aucune réaction. Elles servent d'exemples contrastifs pendant l'entraînement, pour apprendre au modèle à distinguer le vaguement pertinent de l'utile.

Présélection, ou récupération. Première étape du système, qui réduit des millions de publications à quelques centaines de candidates. Distincte du classement, qui les ordonne ensuite.

Transformeur. Famille d'architectures qui traite une suite d'éléments en pondérant l'importance de chacun par rapport aux autres. Dans le cas du fil, la suite est celle de vos interactions passées.

Valeur p. En statistique, probabilité d'observer un écart au moins aussi grand si aucun effet réel n'existait. Une valeur inférieure à 0,001 signifie moins d'une chance sur mille que le résultat soit dû au hasard.

SECTION 14

SOURCES

— Publications de LinkedIn

Hertel, L., Srivastava, S. et al. (2026). *An Industrial-Scale Sequential Recommender for LinkedIn Feed Ranking*. arXiv:2602.12354, 12 février 2026. <https://arxiv.org/abs/2602.12354>

Borisyyuk, F. et al. (2024). *LiRank: Industrial Large Scale Ranking Models at LinkedIn*. Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining, Barcelone, pages 4804 à 4815. arXiv:2402.06859. <https://arxiv.org/abs/2402.06859>

Firooz, H. et al. (2025). *360Brew: A Decoder-only Foundation Model for Personalized Ranking and Recommendation*. arXiv:2501.16450, 27 janvier 2025. Article qualifié de pré-production, retiré depuis d'arXiv. Version consultable : <https://arxiv.labs.arxiv.org/html/2501.16450>

Danchev, H. (2026). *Engineering the next generation of LinkedIn's Feed*. Blog d'ingénierie LinkedIn, 12 mars 2026. Source de l'architecture de récupération, de l'épisode des centiles, des négatifs difficiles et du détail des six actions prédites. <https://www.linkedin.com/blog/engineering/feed/engineering-the-next-generation-of-linkedins-feed>

Jurka, T. (2026). *Updates to The LinkedIn Feed Focusing on Authentic, Relevant Conversations*. Mars 2026. <https://www.linkedin.com/pulse/updates-linkedin-feed-focusing-authentic-relevant-tim-jurka-umwnc/>

LinkedIn Engineering. *Understanding dwell time to improve LinkedIn feed ranking*. <https://www.linkedin.com/blog/engineering/feed/understanding-feed-dwell-time>

— Études indépendantes

Ordinal. *LinkedIn Link Penalty Study*. Plus de 900 000 publications, février 2023 à février 2026. La page a été consultée pour ce document puis est devenue inaccessible (erreur 404 au 27 juillet 2026). Ses résultats restent repris par plusieurs sources secondaires. Site de l'éditeur : <https://www.tryordinal.com/blog>

van der Blom, R. *Algorithm Insights*, édition 2026. 1,3 million de publications, 50 000 créateurs. <https://richardvanderblom.com/>

Saywhat. *State of the Algorithm, Q1 2026*. 397 605 publications. Publication indépendante, sans affiliation à LinkedIn. <https://saywhat.ai/algorithm-webinar/>

AuthoredUp. *Best Performing Content on LinkedIn*. 3 millions de publications, mars 2025 à février 2026. <https://authoredup.com/blog/best-performing-content-on-linkedin>

— Sur la diffusion de l'affirmation 360Brew

Cook, J. (2026). *The LinkedIn Algorithm Changed Again. Here's What's New For 2026*. Forbes, 12 janvier 2026. <https://www.forbes.com/sites/jodiecook/2026/01/12/the-linkedin-algorithm-changed-again-heres-whats-new-for-2026/>

— Pour aller plus loin

PPC Land. *LinkedIn rebuilds its feed from scratch with LLMs and GPU-powered ranking*. <https://ppc.land/linkedin-rebuilds-its-feed-from-scratch-with-llms-and-gpu-powered-ranking/>

Social Media Today. *LinkedIn Updates Its Feed Algorithm*. <https://www.socialmediatoday.com/news/linkedin-updates-its-feed-algorithm/814638/>

Empower Agency. *What LinkedIn's new feed algorithm means for your content strategy*. <https://empower.agency/insights/social-media/what-linkedins-new-feed-algorithm-means-for-your-content-strategy/>



À PROPOS

Fast Growth Advisors est un cabinet de conseil en Message-Market Fit. Nous auditons la clarté du message des entreprises B2B et mesurons leur lisibilité par leurs acheteurs, par les analystes et par les moteurs de réponse. Ce document est publié en accès libre et peut être cité avec mention de la source.

Contact : herve@fast-growth.fr · fast-growth.fr

ALGORITHME LINKEDIN 2026 – CE QUI EST PROUVÉ, CE QUI EST MESURÉ, CE QUI EST INVENTÉ
FAST GROWTH ADVISORS · JUILLET 2026 · DOCUMENT PUBLIC