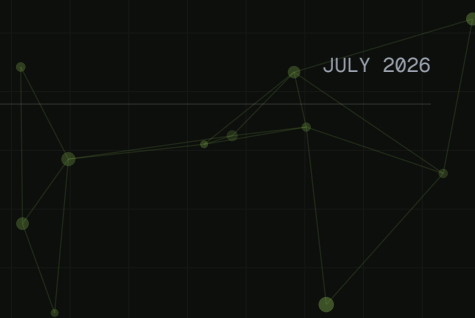


THE LINKEDIN ALGORITHM 2026



What is proven, what is measured,
what is invented

A review of the public record. No advice, no tactics.

Proven

LinkedIn wrote or published it: research paper, engineering blog, official communication.

Measured

An independent study, with published methodology and data volume.

Disputed

Several credible sources reach incompatible conclusions.

Invented

Circulates widely, traces back to nothing. The precision is the only argument.

*This document recommends nothing. It establishes what is known,
and how one can know it.*

NAVIGATION

CONTENTS

This document compares three things: what LinkedIn publishes, what independent studies measure, and what everyone repeats.

01	Summary	03
02	Why a review rather than a guide?	04
03	What does LinkedIn actually publish?	05
04	Is the feed ranked by a generative AI?	06
05	How does the system understand what a post is about?	08
06	What does the ranking model actually optimize for?	10
07	Are external links penalized?	11
08	Do the circulating numbers hold up?	13
09	What nobody knows	14
10	How do you assess a source on this subject?	15
11	Method and limitations	16
12	Frequently asked questions	17
13	Glossary	18
14	Sources	19

A review of the public record. No advice, no tactics.

SECTION 01

SUMMARY

On February 12, 2026, twenty-four LinkedIn engineers published a paper describing the model that ranks the feed for 1.2 billion members. It is openly available. Almost none of the articles explaining "the new algorithm" cite it.

This document compares three things: what LinkedIn publishes, what independent studies measure, and what everyone repeats. The three do not line up.

Four findings come out of it.

The most awkward one first, because it contradicts the most widely repeated claim of the year: in the two documents LinkedIn published in the first quarter of 2026, the feed is not ranked by a generative AI. The company built one and tested it on live traffic. As of its paper, it was not outperforming the model already in service, and it was not in production.

Second finding, simpler: the system optimizes for a small number of specific behaviors, which LinkedIn names one by one in its publications.

The third demands some discomfort. On external links, three serious studies reach incompatible conclusions, one of them the reverse of the other two, so no firm answer is possible as things stand. That leaves the fourth, easier to state: several widely circulated numbers trace back to nothing.

This document recommends nothing. It establishes what is known, and how one can know it.

SECTION 02

WHY A REVIEW RATHER THAN A GUIDE?

There are already hundreds of guides to the LinkedIn algorithm. There is almost nothing on how reliable their claims are.

The contrast is striking. On one side, LinkedIn publishes its work in the research literature, with the details of its models and the results of its experiments. It is free and open to anyone. On the other, what you read on the subject comes from posts commenting on other posts, each adding a layer of certainty to a claim whose origin nobody can trace anymore.

The result: the verifiable and the invented read exactly alike. We set out to separate them. The method is simple: trace every claim back to its primary source, or establish that it has none.

— The four evidence levels used in this document

01

Proven

LinkedIn wrote or published it: research paper, engineering blog, official communication. The highest level of confidence, with one caveat: a company describes what it chooses to describe.

02

Measured

An independent study with a published methodology and data volume. Reliable, but the sample is never representative of the platform as a whole.

03

Disputed

Several credible sources reach incompatible conclusions. That has to be stated, not arbitrated.

04

Invented

Circulates widely, traces back to nothing. The precision of the figure is often the only argument.

SECTION 03

WHAT DOES LINKEDIN ACTUALLY PUBLISH?

Everything rests on three publications. Two of them are almost never cited in what gets written about LinkedIn.

The most important is a research paper from February 2026, *An Industrial-Scale Sequential Recommender for LinkedIn Feed Ranking*, posted on arXiv. Twenty-four engineers sign it. Its abstract states that this model is, at that point, the primary system behind the feed. The paper then describes its objectives, its architecture and its experimental results, at a level of detail no other source comes close to. This is the document you will find almost no trace of in what gets written on the subject.

Next comes a March 2026 post on LinkedIn's engineering blog, by Hristo Danchev, this time about retrieval: how the platform selects the few hundred posts it will then rank. This is the one the trade press picked up, often without reading it.

That leaves two older documents. A 2020 technical note on dwell time, routinely misquoted. And *LiRank*, presented at a research conference in 2024, which describes the previous system.

All four were read in full. That is the only way to see how partial the second-hand accounts of them are.

SECTION 04

IS THE FEED RANKED BY A GENERATIVE AI?

It is the most widely repeated claim of 2026. Both of LinkedIn's publications contradict it.

— What is proven

In the two documents published in the first quarter of 2026, feed ranking is handled by a model the engineering blog calls Generative Recommender and the research paper calls Feed SR. It is a causal-attention sequential transformer, meaning a model that reads your interactions in chronological order without ever looking ahead. It replaced a previous-generation model from a different architectural family. It is not a language model.

Language models are involved, but at two points that are not ranking: retrieval, meaning the selection of candidate posts, and the representation of the profile of the member browsing the feed. Put plainly: a language model decides which posts are candidates. A far smaller model decides the order they appear in.

— Where the misunderstanding comes from

From a sentence by LinkedIn itself, which makes the error fairly understandable. The second paragraph of the March 2026 blog post announces the rollout, in LinkedIn's exact words, of "a new advanced ranking system, powered by LLMs and GPUs". That is the sentence everyone quoted, and it is not false: the system does contain a language model.

The body of the same post says something else, in detail and over several pages. Retrieval does rest on a language model fine-tuned for the task and built as a dual encoder, an arrangement that maps the member and the post separately into two directly comparable representations. Ranking rests on an entirely different model, a sequential transformer that gives each action type its own prediction head. The two stages are distinct, and the post keeps them apart. What went around the web was the opening sentence.

— What LinkedIn says it tried

The paper devotes a full section to the alternative architectures evaluated before the final choice. The team did build a ranker based on a language model. Each post was described as text inside a prompt, and the model answered questions of the form "will this member click?". The paper sets out three obstacles.

The first lies in the nature of the data. A language model works on text, and much of what matters for ranking a post does not express itself in words: reaction counts, recency, closeness to the author. Rendered as sentences, these values lost their usefulness, and the model could not exploit them properly.

Then comes cost. A language model splits text into tokens and is billed by the token. Each post consumed hundreds of them, against two in Feed SR.

And then the decisive obstacle, performance. At the time the paper was written, this ranker was not beating the model already in service on online metrics. The abstract states it without hedging: the authors explain "why Feed-SR provided the best combination of online metrics and production efficiency". It did creditably on out-of-network content and failed on in-network content, the strength of a relationship being hard to put into words.

One clarification here, and it holds for this entire document. LinkedIn describes a result obtained on a given date, not a settled decision. The company has announced no abandonment, and nothing indicates where matters stand since. What is established fits in a single sentence: the system described in the first quarter of 2026 ranks the feed with a sequential transformer, and the language-model route had not, at that point, produced better results.

— So where does the giant-model story come from?

From 360Brew, and the story is worth telling because it shows how a claim drifts without anyone really lying.

360Brew is real. It is a research paper posted on arXiv in January 2025 by a LinkedIn team, describing a 150-billion-parameter model able in principle to handle feed ranking, job recommendations and connection suggestions in a single system. Its abstract calls it, in as many words, a "research pre-production model", that is, an undeployed prototype.

A year later, in January 2026, Forbes published an interview in which a creator with 1.2 million followers states that LinkedIn's latest update is called 360 Brew and that it "now shows your content more accurately to your ICP". Within weeks the claim had become settled fact, and dozens of posts were building detailed tactical recommendations on top of it.

Three things contradict it, all verifiable in minutes. The paper describes itself as pre-production. We found no confirmation of deployment from LinkedIn. And its March 2026 engineering post never mentions the name.

A fourth element shows how far the chain of copying drifted from the source. Several articles state that 360Brew is built on LLaMA 3, when the paper describes a Mixtral architecture. The error concerns something written in plain sight in the text: it tells you the source was not opened before being cited. In twelve months, a lab model became the name commonly given to LinkedIn's algorithm.

SECTION 05

HOW DOES THE SYSTEM UNDERSTAND WHAT A POST IS ABOUT?

Here the engineering blog describes a mechanism nobody picks up, and one that illuminates everything else. Under the heading "Unified retrieval through fine-tuned LLMs", the post explains that LinkedIn turns every post and every member into text, passed to a language model that produces a mathematical representation of it. Matching no longer happens on keywords, but on meaning.

What goes into the text describing a post is listed explicitly: its format, its text, its reaction and view counters, the article's details if there is one, and its author's details, namely name, headline, company and industry. In other words, the author's headline is part of the description of their post.

On the member side, the text gathers the profile, the skills, the career history, the education, and the chronological list of posts the person has interacted with.

— The mistake LinkedIn tells on itself

LinkedIn recounts a mistake, and the anecdote is worth more than a long discussion of the limits of language models. The section "From structured data to effective prompts" tells the story. Engagement counters were passed in raw at first. A post with 12,345 views appeared in the text as "views:12345". The model treated those digits like any other characters, and the outcome alarmed the team: the correlation between a post's popularity and its computed relevance was minus 0.004, in other words nil.

Yet popularity is a strong signal. A post many people find useful stands a good chance of being useful to you too.

The fix was to replace the raw number with its position in the distribution. "views:12345" became "71st percentile", meaning: this post is seen more than 71% of the others. The model could finally learn what "above average" means.

The correlation improved thirtyfold. Retrieval quality rose by 15%.

— What the system learns from your silences

Two training details are worth knowing, because they describe what the machine treats as signal. The first, set out under "Training dual encoders at scale", concerns counter-examples. To learn to tell a vaguely relevant post from a genuinely useful one, the model is trained on posts that were shown to a member without prompting any reaction whatsoever. LinkedIn calls them hard negatives. Adding two per member improved retrieval quality by 3.6%.

The second is more counter-intuitive. A member's history initially included everything shown to them, including what they scrolled past without stopping. The team removed that last set. The model learned better, and training became 2.6 times faster.

The phrasing LinkedIn uses is worth quoting:

the model learns better from what you chose than from everything you were shown.

— The case of the brand-new member

The blog describes a specific scenario it labels "cold-start": a newly registered member, known only by headline and job title, with no history at all. From those two facts alone the language model infers the topics likely to interest them, drawing on what it learned about the world from reading billions of texts. LinkedIn gives an example: someone whose profile says "electrical engineering" but who engages with content on small modular reactors. A keyword system makes no connection. A language model knows those subjects are adjacent.

SECTION 06

WHAT DOES THE RANKING MODEL ACTUALLY OPTIMIZE FOR?

The February 2026 paper answers head-on, and this is probably the most useful fact in the whole file. The model predicts six actions, and the engineering blog names every one of them.

The section "Teaching transformers to recommend" lists them. Three are passive: the click, the scroll-past, and **Long Dwell**, meaning staying on a post beyond a threshold that depends on its type. Three are active: the reaction, the comment and the share. The research paper groups those last three under the name **Contribution**.

The two families are not mixed together. The model handles them separately at prediction time, working from the same reading of your history.

The two metrics LinkedIn foregrounds in its paper, the ones it reports its gains on, are Long Dwell and Contribution.

— What the paper also documents, with numbers

Its memory is bounded: one thousand viewed posts, drawn from roughly a year. Beyond that, nothing. The weight of each decays over time, through two distinct mechanisms. The data used to train the model loses half its influence every sixty days. And within a given member's sequence, the oldest interactions count about half as much as the most recent. A history goes stale, at a measurable rate.

Another signal, less expected: how previous readers break down across dwell-time buckets, from zero to five seconds up to more than sixty. That element alone improves Long Dwell prediction by 2.5 points.

LinkedIn also corrects a bias few people suspect. A post placed at the top of the feed mechanically collects more interactions than another, whatever its quality, simply because it is seen first. The system compensates by weighting each observation by the probability it had of being seen, and by learning an offset specific to each of the first sixty positions. In other words, it tries to measure what a post is really worth, independently of how lucky it was with placement.

Result of the online experiment: 2.10% more time spent, with the largest gains among the most active members and no measurable effect on new ones.

— One consequence for reading the rest of the market

A common claim holds that saving a post, meaning setting it aside to read later, has become the top ranking factor. The paper does not list saves among the objectives it optimizes. That does not prove they count for nothing, but the claim rests on no primary source, whereas the two named objectives are documented.

SECTION 07

ARE EXTERNAL LINKS PENALIZED?

It is the most frequently asked question, and the one where honesty means declining to conclude.

— LinkedIn's position

We found no rule in LinkedIn's public documentation penalizing posts that contain a link. Its documentation describes quality and safety criteria, and reduces the distribution of content judged low-quality or spam. Links do not appear there.

— What the studies measure

They do not agree, and the gap is considerable. A study published by Ordinal, covering more than 900,000 posts across thirty-seven months, reports a 26.5% drop in reach for posts containing a link, with a Mann-Whitney test and a p-value below 0.001, meaning less than a one-in-a-thousand chance that the observed gap is due to chance. It is the only available study to publish that kind of statistical check.

Two caveats come with it, and we flag them because our own framework requires it. Ordinal sells a LinkedIn publishing tool, one function of which automatically moves links into the first comment: the company therefore has a commercial interest in the penalty being real. And on another page of its site, it presents that function as a way to avoid "the 60% external link penalty", which is precisely the unsourced figure we file below among the inventions.

Richard van der Blom's *Algorithm Insights* report, based on 1.3 million posts, reports an 18.8% drop in median reach, meaning the reach of the post sitting in the middle of the distribution, for a link placed in the body of the text. The *State of the Algorithm* report for the first quarter of 2026, published by Saywhat from 397,605 posts, finds the opposite: posts containing several external links perform markedly better there than posts containing none. The authors attribute this to content quality, posts rich in useful resources generating enough engagement signal to compensate.

As for the widely circulated 60% penalty figure, it traces back to no identifiable source.

— The hypothesis that reconciles these results

It follows from what we know about the model. Assuming no link detector penalizes posts at all, a mechanical chain is enough to produce the same outcome.

A reader clicks, and leaves the platform. So they no longer read the post, which collapses Long Dwell. They neither comment nor share, which removes Contribution. The model observes content that does not hold attention and does not prompt a response, and distributes it less.

The link would then not be punished: it would mechanically produce the signals of mediocre content, which amounts to the same thing from the model's point of view but changes everything in the interpretation. There is also a more down-to-earth explanation nobody mentions. Many link posts are hollow, published solely to place a URL. They underperform because they are empty.

That said, this hypothesis is not demonstrated. It is consistent with the whole body of observations, which is not the same thing.

SECTION 08

DO THE CIRCULATING NUMBERS HOLD UP?

We selected four of the most widely shared claims, chosen because they are unverifiable and because each carries a precise figure.

— Engagement rates by dwell-time bucket

You regularly read that engagement goes from 1.2% below three seconds to 15.6% beyond sixty seconds. The buckets exist, the February 2026 paper confirms it. But we found no engagement rate per bucket anywhere in LinkedIn's publications. These numbers are even repeated by sites that, on another of their own pages, note that they cannot be traced.

— The "Depth Score"

Presented since late 2025 as LinkedIn's new criterion. The term appears in none of the LinkedIn publications we consulted, and we found no source attributing it explicitly to the company. The idea behind the word, measuring how long someone stops on a piece of content, is real and has been documented since 2020 under the name dwell time. The name itself appears to come from somewhere else.

— The three percentages circulating in French

They turn up in many French-language articles and posts, always without a reference. Commenting under your own post before anyone else supposedly costs 20% of reach. Replying within the first hour supposedly gains 35%. Commenting elsewhere fifteen to thirty minutes before publishing supposedly adds 21%.

We found none of these three figures, neither in LinkedIn's publications nor in any study whose method and sample are published.

— Engagement rates by format

They are irreconcilable with one another. For carousels, depending on the source, you find 1.44%, 6.60% or 24.42%. None of these tools publishes its calculation method: we do not know what it puts in the denominator, or what sample it works from. Absent that information, comparing them means nothing.

— A word on how these figures spread

A precise number shares better than an honest uncertainty, and it never asks to be checked. That is probably why "minus 18.8% according to a study of 1.3 million posts" travels less well than "minus 60%".

SECTION 09

WHAT NOBODY KNOWS

An honest review must also say what is not known. The exact weight of each factor will never be public. When you read "this criterion counts for X percent", it is an invention.

We also do not know what has happened since March 2026. LinkedIn has kept publishing, but on other subjects: its hiring assistant in June 2026, its internal software-quality tooling. No publication devoted to feed ranking has appeared since the March one, at least none we found. Speaking of the "current" algorithm therefore assumes nothing important has changed in a few months. That is likely. It is not verifiable.

In the detail, several points remain open. Long Dwell thresholds vary by post type, the paper confirms it, but no value is given, which rules out any length rule founded on anything other than your own results. The fate of links placed in comments remains frankly uncertain, some data suggesting degradation and other data not. And the effect of editing a post after publication has never been measured, even though the belief that it destroys reach is everywhere.

The independent studies we consulted all rest on the customers of the tool publishing them, and none states that it works from a random sample drawn across all members. Their results are useful, and skewed from the outset.

On one last point, LinkedIn has to be taken at its word. The company says it audits its models regularly to verify two things: that posts from different authors are treated the same way, and that what some members see resembles what others see. It adds that its models look at professional activity and never at age, gender or origin. These audits are not published.

SECTION 10

HOW DO YOU ASSESS A SOURCE ON THIS SUBJECT?

Four questions are enough to filter out most of the noise, and none of them calls for technical skill.

— Does the source cite a verifiable primary publication?

A link to arXiv or to LinkedIn's engineering blog, not to another blog post. A source that never cites arXiv has not read LinkedIn's work.

— Does it give its data volume and its time window?

Without a sample and a time frame, a percentage is worth nothing.

— Does it sell something that its conclusion serves?

A carousel-tool vendor concluding that carousels perform best is not a neutral observer.

— Is the figure too precise for its origin?

A value drawn from a named, dated study can be precise. A general claim carrying a round percentage, with no source, never legitimately is.

SECTION 11

METHOD AND LIMITATIONS

This document was established in July 2026 and covers the public state of knowledge at that date.

A word on the vocabulary used. When we write that a piece of information does not exist, read it as: we did not find it in the sources listed at the end of this document. Proving the absence of a thing is impossible. Reporting that it cannot be found anywhere is verifiable by anyone who retraces the path.

LinkedIn's publications describe systems and experimental results, never the exact weight of each criterion. What appears in them is therefore accurate and necessarily incomplete, and we assert nothing beyond what they write. This document will age, too. What it describes dates from February and March 2026, the platform moves fast, and the principles will hold up better than the numbers.

We have covered neither advertising, nor company pages, nor job recommendations. Those are different systems, and they deserve separate examination.

SECTION 12

FREQUENTLY ASKED QUESTIONS

Does LinkedIn really publish how its algorithm works?

In part, yes. The company publishes the details of its models, its metrics and the results of its experiments on real users, in the research literature and on its engineering blog. It does not publish the weights of its models. What is described is therefore real, but necessarily incomplete.

Why do so few articles cite these publications?

They are technical and long, so poorly suited to the short format. It is quicker to comment on a blog post already commented on elsewhere. The price of that convenience: after a few rounds of copying, nobody knows where the information came from.

Are the independent studies reliable?

They are useful, and skewed from the outset. The ones we consulted cover the customers of the tool publishing them, and none states that it works from a random draw across all members. When they compare two situations within the same group, they can be trusted. When they announce an absolute figure, far less so.

What should you retain if you retain only one thing?

That the ranking model predicts six named actions, grouped into two families, and that any claim which acts on none of them deserves to be questioned.

Does the author's headline influence how their posts are distributed?

The engineering blog lists what makes up the description of a post sent to the model: its text, its format, its counters, and its author's details, including headline, company and industry. The headline is therefore part of what situates a piece of content. LinkedIn does not quantify that effect.

Does this document give recommendations?

No, and that is deliberate. It establishes what is known and how one can know it. What to do with that depends on each situation.

SECTION 13

GLOSSARY

The terms below recur in LinkedIn's publications. They are given in the exact sense the company gives them, not in a general sense.

AUC. Standard measure of the quality of a prediction, ranging from 0.5 to 1. An AUC of 0.5 is equivalent to chance. Gains reported in this field are counted in fractions of a point, which is why a two-point gain is presented as significant.

Percentile. The position of a value within a distribution. A post in the 71st percentile of views is seen more than 71% of the others. LinkedIn replaced raw counters with percentiles to make them usable by its models.

Contribution. Term used by LinkedIn to designate the set formed by the reaction, the comment and the share. It is one of the two objectives on which the company reports its gains.

Dual encoder. An architecture in which one model processes two objects separately, here a member and a post, to produce two representations that are directly comparable with each other.

Feed SR, or Generative Recommender. The ranking model in production. Two names for the same system: the first comes from the research paper, the second from the engineering blog.

Long Dwell. Staying on a post beyond a time threshold, a threshold that varies by post type. The second objective on which LinkedIn reports its gains. Not to be confused with dwell time, which is the raw measure of which Long Dwell is a threshold.

Language model. A model trained on very large volumes of text, able to establish connections of meaning that no explicit rule taught it. At LinkedIn it serves retrieval of posts, not their ranking.

Hard negatives. Posts actually displayed to a member that prompted no reaction. They serve as contrastive examples during training, to teach the model to tell the vaguely relevant from the useful.

Retrieval. The first stage of the system, which narrows millions of posts down to a few hundred candidates. Distinct from ranking, which then orders them.

Transformer. A family of architectures that processes a sequence of elements by weighting the importance of each relative to the others. In the case of the feed, the sequence is the one formed by your past interactions.

p-value. In statistics, the probability of observing a gap at least as large if no real effect existed. A value below 0.001 means less than a one-in-a-thousand chance that the result is due to chance.

SECTION 14

SOURCES

— LinkedIn publications

Hertel, L., Srivastava, S. et al. (2026). *An Industrial-Scale Sequential Recommender for LinkedIn Feed Ranking*. arXiv:2602.12354, February 12, 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, Barcelona, pages 4804 to 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, January 27, 2025. Described as pre-production, since withdrawn from arXiv. Readable version: <https://arxiv.labs.arxiv.org/html/2501.16450>

Danchev, H. (2026). *Engineering the next generation of LinkedIn's Feed*. LinkedIn engineering blog, March 12, 2026. Source for the retrieval architecture, the percentile episode, the hard negatives and the detail of the six predicted actions. <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*. March 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>

— Independent studies

Ordinal. *LinkedIn Link Penalty Study*. More than 900,000 posts, February 2023 to February 2026. The page was consulted for this document and has since become unreachable (404 as of July 27, 2026). Its results are still relayed by several secondary sources. Publisher site: <https://www.tryordinal.com/blog>

van der Blom, R. *Algorithm Insights*, 2026 edition. 1.3 million posts, 50,000 creators. <https://richardvanderblom.com/>

Saywhat. *State of the Algorithm, Q1 2026*. 397,605 posts. Independent publication, no affiliation with LinkedIn. <https://saywhat.ai/algorithm-webinar/>

AuthoredUp. *Best Performing Content on LinkedIn*. 3 million posts, March 2025 to February 2026. <https://authoredup.com/blog/best-performing-content-on-linkedin>

— On the spread of the 360Brew claim

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

— Further reading

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/>



ABOUT

Fast Growth Advisors is a consultancy specializing in Message-Market Fit. We audit the clarity of B2B companies' messaging and measure how legible it is to their buyers, to analysts and to answer engines. This document is published openly and may be cited with attribution.

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

THE LINKEDIN ALGORITHM 2026 – WHAT IS PROVEN, WHAT IS MEASURED, WHAT IS INVENTED
FAST GROWTH ADVISORS · JULY 2026 · PUBLIC DOCUMENT