

How the LinkedIn Algorithm Works: Two-Pass Retrieval, Multi-Objective Ranking, and the Determinants of Reach

Ruchir Jajoo¹ and Husain

Ruchir Jajoo, Founder, **Identity Labs**, India

Husain, Technical Lead, **Identity Labs**, India; IIT Delhi

¹Corresponding author: ruchir@identity-labs.com

LinkedIn's home feed retrieves tens of thousands of candidate updates per request and returns an ordered feed of a few dozen. The system performing this reduction has never been released as source code, and none of its ranking coefficients have been published. It is nonetheless unusually well documented: across a decade of engineering publications and six peer-reviewed papers, LinkedIn has described its retrieval architecture, its objective function, and its model families in specific technical detail. We reconstruct the pipeline from that literature and organise it into five stages: signal collection, quality filtering, first-pass candidate generation, second-pass multi-objective ranking, and feed construction. We formalise the published ranking objective as a weighted combination of passive-consumption, active-contribution, and creator-side utility terms, and document the model families that produce each term — a two-tower multi-task network over XGBoost leaf embeddings, a P(skip) dwell-time model, and the Residual DCN architecture reported in production as of 2024. We then grade every quantitative claim in circulation by evidential tier, separating what LinkedIn has documented from what practitioners have measured externally and what is merely inferred. The central finding is an asymmetry: the architecture is specified to a level permitting reimplementability, while the parameters are not merely unpublished but continuously retuned online, so that no fixed constant exists to disclose. Widely circulated practitioner figures — the claim that comments are weighted fifteen times a reaction, for instance — vary by roughly an order of magnitude across sources and rest on uncontrolled observational data. We delimit what can and cannot be concluded about author-side determinants of reach on this basis.

Recommender systems | Feed ranking | Multi-objective optimisation | Multi-task learning | Algorithmic transparency

Introduction

A member opening LinkedIn initiates a retrieval problem of substantial scale. The candidate pool comprises every recent update authored by their first- and second-degree connections, by the pages and people they follow, and by authors outside their network whom the system judges relevant — on the order of tens of thousands of content records per request, by LinkedIn's own account of its retrieval tier [1]. From this pool the feed returns a few dozen ordered updates. The selection ratio is comparable to that of other population-scale feed systems, and the mechanism performing the selection is correspondingly consequential: it determines which professional content is seen, by whom, and how often.

This paper reconstructs that mechanism. Our situation differs materially from the equivalent exercise on X, whose recommendation source was published in March 2023 and whose ranking weights can therefore be read directly from a configuration file. LinkedIn has released no source and no coefficients. What it has released is a decade of technical writing about its own system: engineering-blog descriptions of the retrieval tier [1], the candidate-selection objective [2], and the ranking model [3]; and peer-reviewed papers on feed personalisation [4], dwell-time modelling [5], creator-side utility [6], online parameter selection [7], and the production ranking stack [8].

The result is an inversion of the usual transparency trade-off. Where the X release gives exact constants embedded in

an incompletely described system, the LinkedIn literature gives a well-described system with no constants at all. Each admits a different class of conclusion, and conflating them produces the confident numerical claims about LinkedIn's feed that circulate widely and survive scrutiny poorly.

Our contribution is threefold. First, we consolidate the published architecture into a single account in which each stage's inputs, transformations, and model families are identifiable (Fig. 1). Second, we state the ranking objective in the form LinkedIn has published it, and decompose it into the predicted quantities that enter it. Third — and this is the methodological core of the paper — we grade every quantitative claim by evidential tier, and show that the figures most often cited as facts about LinkedIn's ranking are external estimates whose values disagree across sources by roughly an order of magnitude.

Evidential convention

Throughout, we mark quantitative and mechanistic claims according to their source. This convention is applied uniformly, including to claims we consider well founded.

- D Documented.** Stated by LinkedIn in an engineering publication or a peer-reviewed paper with LinkedIn authorship. Reproducible by reading the cited source.
- R Reported.** Measured by external practitioners from observed post performance. Not confirmed by LinkedIn; derived from uncontrolled observational samples.
- I Inferred.** Follows by reasoning from documented mechanisms, but is not itself stated in any source.

System Overview

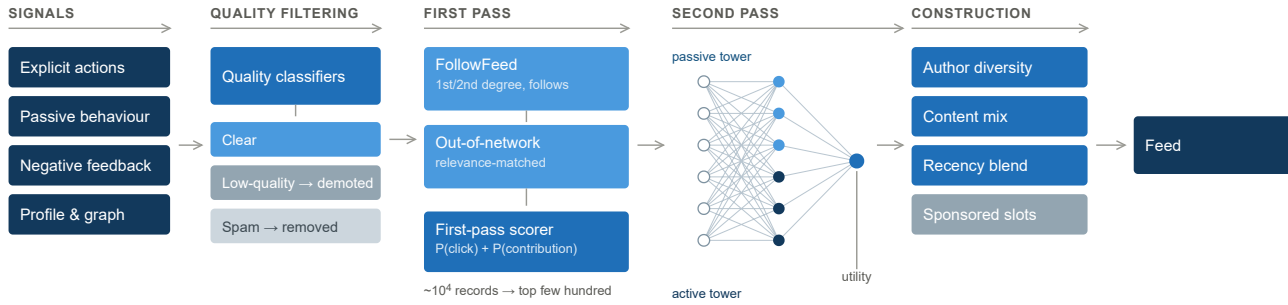


Fig. 1. Architecture of the LinkedIn home feed pipeline. Member and content signals (left) are grouped into explicit actions, passive behaviour, negative feedback, and profile and graph attributes, and are converted into the features every downstream stage consumes. Quality classifiers partition candidates into clear, demoted, and removed sets before ranking begins. The first pass retrieves from two sources — FollowFeed, which indexes updates from the member’s network, and out-of-network retrieval — and scores the pool with a lightweight ensemble optimising a joint click-and-contribution objective, reducing tens of thousands of records to a few hundred. The second pass scores that set with a multi-task network whose passive-consumption and active-contribution towers share an input representation and converge on a single utility. Construction applies diversity, mixing, and recency adjustments and inserts sponsored slots before the feed is served. Stage names follow LinkedIn’s own terminology [1]–[3]; no ranking coefficients are published for any stage.

The pipeline decomposes into five stages, of which the third and fourth carry nearly all of the discriminative work. Figure 1 shows the arrangement.

Signal collection and feature engineering converts raw interaction events into the feature representation every later stage consumes. *Quality filtering* applies classifiers that remove or demote spam and low-quality content before any expensive scoring occurs. *First-pass candidate generation* retrieves a few hundred plausible candidates from a pool of tens of thousands using a deliberately cheap model. *Second-pass ranking* scores that reduced set with a considerably more expensive multi-task network. *Feed construction* applies diversity and business constraints and interleaves sponsored content.

Two structural properties of this arrangement matter for everything that follows. First, the two passes optimise different objectives against different budgets: the first pass is evaluated on *recall* against the second pass’s preferences [2], not on final ranking quality. Content can therefore fail to reach a feed in two distinct ways — by never being retrieved, or by being retrieved and then ranked poorly — and these have different remedies. Second, the objective is explicitly multi-stakeholder: the published utility function balances the viewer’s interest against value accruing to the content’s author [3][6], a design choice with no counterpart in a purely engagement-maximising system.

Signals and Features

Every interaction on the platform is logged and converted into features. LinkedIn groups the raw signals into four families, and the ranking literature refers to features derived from each [4].

Table 1. Signal families and the feature classes derived from them.

Signal family	Constituents
Explicit action	Reactions, comments, reshares, clicks
Passive behaviour	Dwell time, scroll velocity, expansion of truncated text
Negative feedback	Hide, unfollow, report, “see fewer posts like this”

Member attributes	Industry, skills, connections, followed entities, group membership
Content attributes	Author, format, topic, age, in-network engagement to date

Feature classes derived from these: viewer–content affinity (prior interaction between viewer and author, topic match against the viewer’s interests), content quality, temporal decay, and social proof (engagement already received from the viewer’s connections).

The distinction between explicit and passive signals is more consequential than it first appears. Reactions and comments are sparse, self-selected, and observable to the author; dwell time is dense, involuntary, and invisible. Because dwell time is recorded for every impression rather than only for the small fraction that produce an action, it supplies orders of magnitude more training signal, and LinkedIn has described building ranking objectives on it directly [5]. An author optimising only for visible engagement is optimising against a minority of the signal the system actually receives.

Quality Filtering

Before ranking, automated classifiers partition candidates into three sets: content that is clear, content judged low-quality, and content judged spam. Clear content proceeds; low-quality content is demoted rather than removed; spam is withheld. Borderline cases are escalated to human review.

Two properties of this stage govern its practical effect. It operates *before* ranking, so a demotion applied here cannot be recovered by strong engagement later — the content simply competes from a lower base, or never enters the candidate pool at all. And it is applied without notification: an author whose content is demoted at this stage observes only reduced reach, with no signal distinguishing that cause from ordinary ranking outcomes. The published triggers are generic — engagement bait, excessive tagging, abnormal posting frequency — and LinkedIn has not published the classifiers, their thresholds, or their false-positive rates.

First-Pass Candidate Generation

The retrieval tier is the stage LinkedIn has described in the greatest architectural detail, and it is where the largest reduction in candidate volume occurs (Fig. 2).

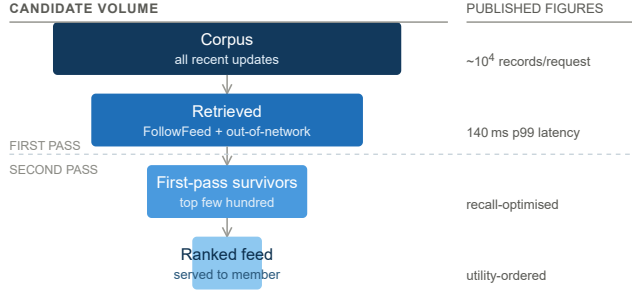


Fig. 2. Candidate volume across the two passes. The retrieval tier returns on the order of ten thousand content records per request and the first-pass scorer reduces this to a few hundred, which the second pass then orders. The first pass is tuned for recall against the second pass’s preferences rather than for final ranking quality, so its errors are unrecoverable: content not retrieved here cannot be ranked at all. Volumes and the latency figure are as published [1][2].

FollowFeed

In-network retrieval is served by FollowFeed, which LinkedIn has described as a fan-out-on-read system organised around *timelines* — per-entity, reverse-chronological lists of content records, keyed by the tuple of entity identifier and content type [1]. Records persist in RocksDB as linked lists of serialised blobs, over-partitioned into 720 partitions [D] so the cluster can be rebalanced without re-sharding.

A service termed the Broker receives each feed request, translates the requesting member’s connections and follows into timeline keys, maps those keys onto partitions, and fans out to the index nodes. Each node performs a batch retrieval and scores its own records in parallel; the Broker then aggregates, deduplicates, and applies diversity filtering before returning the ranked set. Retrieval spans tens of thousands of content records per request [D], and LinkedIn reports p99 latency for the mobile feed of roughly 140 ms [D] — achieved in part by generating specialised Java scoring code, reported at 50µs p99 per record and fifteen times faster than the general scoring library it replaced [D].

The architectural consequence for authors is that in-network eligibility is a graph property, not a quality property. A post enters the retrieval set for a given member if and only if its author occupies a timeline that member’s request expands. No amount of content quality substitutes for that adjacency.

The contribution objective

The first-pass scorer originally predicted click probability alone. LinkedIn subsequently rebuilt it to predict a second quantity it terms *contribution probability* — the likelihood that a member will “share, comment, or react to a particular feed update” [2]. Both objectives are combined by differential weighting inside a single XGBoost tree ensemble, replacing the separate logistic-regression models used previously and halving the parameter space [D].

Two further details of that rebuild are worth recording. The success metric is *recall* — the overlap between the first pass’s top-K and the set the second-pass ranker would have preferred [D] — which confirms that the retrieval tier is explicitly tuned to avoid discarding content the ranker would have rated highly. And candidate selection was

adjusted for *freshness* on the creator’s behalf: LinkedIn observed that viewers are relatively insensitive to update age while creators value prompt feedback, and biased retrieval accordingly [D].

That last adjustment is the earliest point in the pipeline at which creator-side utility, rather than viewer-side relevance, changes what is retrieved. It is also the mechanism underlying the widely reported importance of engagement received shortly after publication [R]: a documented freshness bias in retrieval is consistent with such an effect, though LinkedIn has published neither the size of the bias nor the window over which it decays.

Second-Pass Ranking

The second pass scores the surviving few hundred candidates and produces the order in which they are served. This is the stage LinkedIn has described most fully in its objective and least fully in its parameters.

The objective function

LinkedIn states the ranking utility as a combination of three families of objective [3]: passive-consumption terms F_p covering clicks and dwell time; active-contribution terms F_a covering comments, reshares and reactions; and a residual family F_o covering objectives falling into neither category, principally creator-side feedback. These are combined with tuning parameters the published description denotes α and λ :

$$U = F_p + \alpha F_a + \lambda F_o \quad [1]$$

LinkedIn describes α and λ as balancing ecosystem health between engaged contributors and passive readers [D]. Their values are not published. Figure 3 decomposes each term into the quantities the model actually predicts.

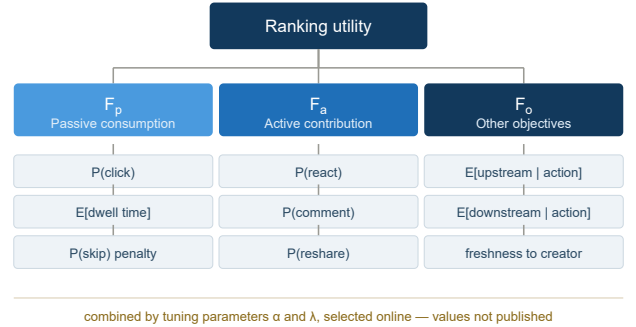


Fig. 3. Decomposition of the ranking utility. Each of the three published objective families resolves into quantities estimated per viewer–content pair. The structure is documented; the coefficients combining the terms are not, and are selected online rather than fixed as constants [3][7].

Predicted quantities

For each viewer–content pair the model estimates three classes of quantity. $P(\text{action})$ is the probability the viewer takes a given action — clicking, reacting, commenting, resharing. $E[\text{downstream} \mid \text{action}]$ is the expected onward effect within the viewer’s own network should they act: a comment can surface the content to the commenter’s connections, so the term quantifies the propagation an action induces. $E[\text{upstream} \mid \text{action}]$ is the expected value

to the author — the feedback and validation the action delivers, and the consequent likelihood of further authoring.

The upstream term is not incidental. LinkedIn devoted a KDD paper to it under the name *feedback shaping*, framing the feed as an instrument for nurturing content creation and reporting that reshaping the distribution of feedback to creators improved content creation significantly without degrading the consumer experience [6] [D](#). A feed optimising Equation 1 is therefore not maximising viewer engagement alone; it trades some viewer relevance for creator retention, deliberately and by design.

Dwell time and the skip model

Dwell time enters ranking through an explicitly negative formulation. LinkedIn measures two quantities: on-feed dwell, which begins when at least half of an update is visible, and post-click dwell, the time spent after opening it [5]. It then identifies a threshold T_{skip} below which a member is judged to have skipped rather than read, and trains a model of

$$P(\text{skip}) = P(\text{dwell} < T_{\text{skip}}) \quad [2]$$

whose output reduces the update’s score proportionally [D](#). LinkedIn reports that T_{skip} proved stable across heterogeneous content types — text, image, video — suggesting members apply a consistent skip decision regardless of format, and that incorporating the model produced fewer skipped updates, more clicks and viral actions, and up to a 10% improvement in model ROC [D](#).

The value of T_{skip} is not published. This matters, because the figure most often quoted in practitioner guidance — that content must hold attention for fifteen seconds [R](#) — is an external estimate of a threshold LinkedIn has stated exists without stating its magnitude. The mechanism is documented; the number attached to it is not.

Model architecture

Rather than training one model per objective, LinkedIn trains a single multi-task network with two towers — one for passive-consumption objectives, one for active-contribution objectives — sharing an input representation built by embedding XGBoost tree leaves as categorical features [3] [D](#). Encoding those leaves densely rather than sparsely reduced serialisation cost to a third of its previous value [D](#).

The production stack has since been described more fully under the name LiRank [8], which reports a residual variant of DCNv2 with attention and residual connections, dense gating, and transformer components unified in a single model, together with isotonic calibration and quantisation for serving. The reported feed result is an increase of 0.5% in member sessions [D](#) — a figure worth holding beside the far larger effects routinely attributed to individual authoring tactics.

The Weighting System

We now address the question practitioners most often ask of this system — which signals count for how much — and the reason it cannot be answered in the form it is usually posed.

Why no constants exist

On X, the analogous question has a literal answer: a configuration file assigns a predicted reply that the author engages with a weight of 75.0 and a predicted report a weight of −369.0, and those numbers can be read directly. No equivalent artefact exists for LinkedIn, and the reason is architectural rather than merely a matter of disclosure policy.

LinkedIn has published a dedicated paper on how the coefficients in Equation 1 are chosen [7]. The parameters are not fixed constants compiled into a ranking service; they are selected online, by a search procedure operating against live business metrics, and updated as those metrics move [D](#). There is consequently no stable weight vector to publish. Asking for LinkedIn’s comment weight is closer to asking for the current value of a control variable than for a documented constant — it has a value at any instant, that value is not the same next month, and it differs across the member segments the optimiser treats separately.

This is the LinkedIn counterpart to the caveat that applies to X’s remotely configurable parameters, but it is considerably stronger. On X, remote configuration can override a published default; on LinkedIn, there is no default to override.

What can be said about ordering

Magnitudes are unavailable, but ordering is partly documented. LinkedIn states that it added contribution probability to candidate selection precisely because click prediction alone under-served conversation [2], that active contribution enters the objective under its own tuning parameter [3], and that predicted skips reduce scores [5]. Together these establish the *direction* of the system’s preferences without establishing their size (Fig. 4).

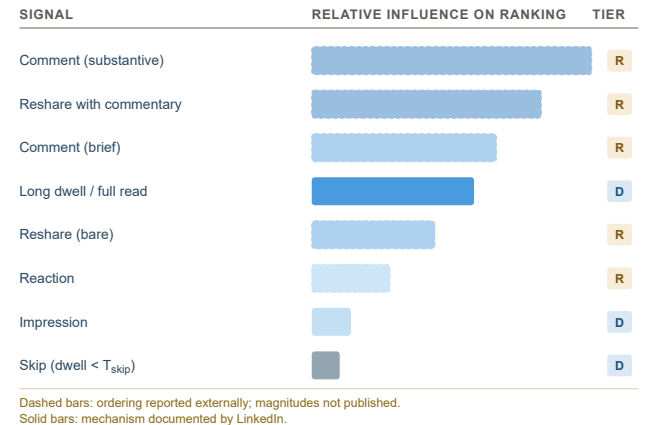


Fig. 4. Signals ordered by ranking influence, graded by evidence. Solid bars denote mechanisms LinkedIn has documented. Dashed bars denote an ordering asserted by external measurement, where the relative magnitudes shown are illustrative rather than published. The rank order is broadly agreed across sources; the spacing between rungs is not, and should not be read quantitatively.

Table 2. Quantitative claims in common circulation, with evidential status.

Claim	Tier	Status
Retrieval spans $\sim 10^4$ records per request	D	Stated [1]
Feed p99 latency ~ 140 ms	D	Stated [1]
Skip threshold T_{skip} exists and is format-invariant	D	Stated [5]
T_{skip} equals 15 seconds	R	Not published; external estimate
Comments weighted $15\times$ reactions	R	Estimates range $\sim 2\times\text{--}15\times$
Creator replies add $\sim 20\%$ distribution	R	Not published; single-source
Complete profiles receive $17\times$ more views	R	Profile-surface claim, not feed ranking
$\sim 400\text{--}450$ views per follower gained	R	Sample-specific; varies by niche
External links suppress reach	I	Consistent with dwell objective; not stated
Posts should be spaced 48 h apart	I	Consistent with diversity rules; not stated

Tiers as defined in the introduction. Rows marked **R** originate in practitioner studies of observed post performance; rows marked **I** are inferences we or others have drawn from documented mechanisms, and no source states them directly.

The dispersion problem

The most-cited figure in circulation holds that a comment advances reach roughly fifteen times as much as a reaction. Other analyses of the same question report values nearer twice, once comment quality is controlled for. A quantity whose published estimates differ by a factor of seven is not a measured constant; it is an artefact of differing samples, differing definitions of reach, and differing treatment of confounders.

The principal confounder is straightforward and, in observational data, not separable. Content that attracts many substantive comments is typically content that is good, timely, and well matched to its audience — and those same properties independently raise every term in Equation 1. Observing that heavily-commented posts travel further does not establish that the comments caused the travel; the comments and the reach may be joint consequences of content quality. Distinguishing the two requires an intervention on comments holding content fixed, which no external party can perform and no published LinkedIn experiment reports.

We therefore state the defensible version: comments are documented to enter the objective under their own term and to have been added deliberately to candidate selection, so their influence is real and directionally established. Its magnitude relative to a reaction is not known outside LinkedIn, and the specific multipliers in circulation should be treated as folklore of uncertain provenance rather than as parameters.

Worked Example

To make the pipeline concrete, we trace a single post through all five stages. The post is a real one: a comparison of cold email against LinkedIn outbound, published as a multi-image carousel by an author in the sales-technology sector, opening with the line “After running outbound for

120+ companies” and closing with a direct question to the reader. It received 135 comments. We reconstruct why, stage by stage. The reconstruction is inferential throughout — we observe the outcome and the documented mechanisms, not the system’s internal state.

TRACING ONE POST THROUGH THE PIPELINE

1. Quality filtering. Structured professional content with a specific quantitative claim and no tag-stuffing or engagement-bait phrasing. Clears the classifiers; no demotion applied.

2. First-pass retrieval. Enters FollowFeed timelines for the author’s first- and second-degree network. The scorer evaluates click probability jointly with contribution probability **D**; a closing question directed at the reader raises the latter. Recency bias favours the post while it is new **D**.

3. Second-pass ranking, passive tower. A multi-image carousel requires sequential swiping to consume. Predicted dwell rises accordingly and predicted $P(\text{skip})$ falls, removing the score reduction that Equation 2 would otherwise apply **I**.

4. Second-pass ranking, active tower. A binary question — one approach or the other — gives readers a position to take and something to say. $P(\text{comment})$ rises; the audience is professionally homogeneous, so $E[\text{downstream} \mid \text{comment}]$ is high, each commenter’s network resembling the author’s target audience **I**.

5. Construction. Author diversity constrains how many of this author’s posts occupy adjacent slots, but does not limit the reach of a single post.

Read against Fig. 3, the post scores through several terms at once rather than one: the carousel format acts on the passive tower, the closing question on the active tower, and the audience homogeneity on the downstream term. This is the pattern the documented architecture rewards. It is also the pattern that makes causal attribution impossible from outside — we cannot say which term contributed what, only that the content was constructed to score on several simultaneously.

One further observation. Nothing in this trace depends on the post being viral in the general sense. Its audience was narrow and professionally specific. LinkedIn’s objective contains no term rewarding breadth as such, and the downstream term rewards propagation into *relevant* networks rather than large ones. Content highly specific to a professional community can therefore propagate efficiently within it without ever becoming broadly popular **I**.

Determinants of Reach

We now invert the analysis. Given the structure above, at which points can an author’s behaviour alter the distribution their content receives? Four such points exist, corresponding to the first four stages. We treat each in turn and mark the evidential basis of every claim.

At the filtering level

The objective here is entirely negative: avoid classification as low-quality. Because this stage precedes ranking, a demotion applied here cannot be offset by engagement afterwards, which makes it the highest-leverage stage in the pipeline and the only one offering no upside — clearing the filters confers no advantage, merely the absence of a penalty.

LinkedIn names engagement bait, excessive tagging, and abnormal posting frequency as triggers, without publishing thresholds [D](#).

At the candidate-generation level

Two documented properties govern this stage. Retrieval is a graph operation: content is eligible for a member’s feed if its author occupies a timeline that member’s request expands [D](#), so the composition of an author’s network sets a ceiling on in-network reach that content quality cannot raise. And the scorer optimises contribution probability alongside click probability [D](#), so content predicted to draw a response is favoured at retrieval, before ranking begins.

Freshness bias applies here on the creator’s behalf [D](#). The widely reported importance of engagement within the first hour [R](#) is consistent with that documented bias, but LinkedIn has published neither the magnitude of the bias nor the shape of its decay, and the specific windows quoted in practitioner guidance — thirty minutes, sixty minutes, ninety minutes — have no published basis.

At the ranking level

Equation 1 identifies three channels, and they are separable. The passive channel is governed by predicted dwell and the skip penalty of Equation 2 [D](#): formats requiring sequential consumption — carousels, documents, native video — mechanically extend the time an engaged reader spends, and text that resolves its opening premise slowly does the same. The active channel is governed by predicted comment, reshare, and reaction [D](#), of which comments are the term LinkedIn explicitly added [D](#). The creator channel is governed by upstream value [D](#), which an author influences only indirectly.

The downstream term deserves separate note, because it is the one most often misread. It rewards an action’s propagation into the acting member’s network [D](#). Engagement from a member whose connections resemble the author’s intended audience is therefore worth more than identical engagement from a member whose connections do not — a property that makes reciprocal-engagement arrangements among unrelated accounts substantially less effective than their participants assume [I](#).

At the construction level

Author diversity limits how many posts from one author occupy adjacent feed positions [D](#). The practical implication is that posting frequency and per-post reach trade against one another within a fixed audience, though the published material does not specify the exchange rate, and the commonly cited 48-hour spacing rule [I](#) is an inference from the existence of diversity constraints rather than a stated parameter.

Summary of determinants

Grouped by evidential tier rather than by stage:

Documented mechanisms.

- Contribution probability enters candidate selection as an explicit objective, so predicted response affects retrieval, not merely ranking [\[2\]](#).
- Predicted skips reduce scores, so formats and structures sustaining attention are rewarded through a documented term [\[5\]](#).

- Network adjacency determines in-network eligibility; retrieval is a graph traversal [\[1\]](#).
- Creator-side utility enters the objective independently of viewer relevance [\[6\]](#).
- Retrieval is biased toward fresh content on the creator’s behalf [\[2\]](#).

Reported but unverified.

- The multiplier by which a comment exceeds a reaction; estimates span roughly 2× to 15×.
- The value of T_{skip} , commonly quoted as fifteen seconds.
- The distribution gain from an author replying to commenters, commonly quoted near one fifth.
- The width of the early-engagement window.

Inferred, not stated.

- Suppression of posts carrying external links. Consistent with a dwell-based objective, since a click leaving the platform truncates measurable dwell, but not stated in any LinkedIn source.
- Optimal spacing between posts.
- The reduced value of engagement from audiences unrelated to the content’s subject.

Limitations

Several limitations bound what can be concluded here, and we state them explicitly.

No source has been released. Every architectural statement in this paper derives from LinkedIn’s own descriptions of its systems rather than from inspection of running code. Those descriptions are written for technical communication, not for reproduction, and they omit what their authors judged uninteresting. We have no means of verifying that the deployed system matches them, nor of detecting components that were never described at all.

No parameters exist to disclose. This is a stronger limitation than non-publication. Because the coefficients of Equation 1 are selected online against live metrics [\[7\]](#), there is no stable weight vector that LinkedIn could publish even under a policy of full transparency. Statements of the form “LinkedIn weights x at n ” are therefore not merely unverified but ill-formed. The most that can be established externally is ordering, and only where a mechanism is documented.

External estimates rest on observational data. The practitioner literature supplying every [R](#)-tier figure in [Table 2](#) measures the association between post attributes and observed reach across samples of posts. It cannot separate the effect of an attribute from the effect of the content quality correlated with it, because the necessary intervention — varying comments while holding content constant — is unavailable to any external party. The resulting estimates are upper bounds on causal effect at best, and the dispersion among them is itself evidence of how much unmodelled variation remains.

The literature lags production. Our sources span 2015 to 2024. The retrieval architecture we describe was published in 2016 and the ranking stack in 2024; the intervening components are documented at different dates and need not describe a single coherent system as it stands today. LiRank

[8] supersedes parts of the multi-task description [3], and further changes since are not observable from outside.

Scope. This analysis addresses the home feed. Search, notifications, jobs recommendation, and advertising ranking share infrastructure and some model families but are governed by distinct objectives and lie outside our scope. Claims about profile-surface behaviour, including the profile-completeness figure in Table 2, concern a different system entirely and are included only because they circulate as though they were feed-ranking claims.

Conclusion

LinkedIn’s home feed reduces tens of thousands of candidate records to an ordered feed of a few dozen through a five-stage pipeline whose architecture is, on the evidence of a decade of the company’s own technical publications, substantially public. Retrieval is a graph traversal over per-entity timelines, tuned for recall against the ranker’s preferences. Ranking is a multi-task network estimating engagement probabilities, dwell, and network propagation, combined into a single utility that includes an explicit term for value accruing to the author rather than the viewer.

Read as a statement of the system’s objective, the published record is clear in its direction. Conversation was added to candidate selection because click prediction alone under-served it. Attention is modelled negatively, through the probability that a member skips rather than the

probability that they engage. Creator retention enters the objective on its own terms, with a dedicated paper behind it. These are design commitments, and they are legible.

What is not legible is any magnitude. The coefficients balancing these objectives are selected online against live business metrics, which means that no fixed constant exists to be published, and that the confident numerical claims circulating about LinkedIn’s ranking — the fifteenfold comment multiplier foremost among them — describe a quantity that does not have a single value. Our comparison with the X release is instructive precisely because the two failures are opposite: X published exact constants for a system whose learned component it withheld, while LinkedIn has described a system thoroughly while withholding, and arguably not possessing, the constants. Neither release permits an external party to reproduce a ranking decision.

For authors, the practical consequence is that the reliable guidance is structural rather than numerical. Network composition bounds retrieval; formats that sustain attention avoid a documented penalty; content predicted to draw a response is favoured before ranking begins; and engagement propagating into a relevant network is worth more than engagement that does not. These follow from documented mechanisms. The multipliers do not, and should be retired from serious discussion of the system until someone inside it publishes a number.

References

1. LinkedIn Engineering (2016) FollowFeed: LinkedIn’s feed made faster and smarter. Engineering blog. [linkedin.com/blog/engineering/feed/followfeed-linkedin-s-feed-made-faster-and-smarter](https://engineering.linkedin.com/blog/engineering/feed/followfeed-linkedin-s-feed-made-faster-and-smarter)
2. LinkedIn Engineering (2020) Community-focused feed optimization. Engineering blog. [linkedin.com/blog/engineering/feed/community-focused-feed-optimization](https://engineering.linkedin.com/blog/engineering/feed/community-focused-feed-optimization)
3. LinkedIn Engineering (2021) Homepage feed multi-task learning using TensorFlow. Engineering blog. [linkedin.com/blog/engineering/feed/homepage-feed-multi-task-learning-using-tensorflow](https://engineering.linkedin.com/blog/engineering/feed/homepage-feed-multi-task-learning-using-tensorflow)
4. Agarwal D, Chen B-C, He Q, *et al.* (2015) Personalizing LinkedIn feed. *Proc. 21st ACM SIGKDD Conf. on Knowledge Discovery and Data Mining*, pp 1651–1660. doi.org/10.1145/2783258.2788614
5. LinkedIn Engineering (2020) Understanding dwell time to improve LinkedIn feed ranking. Engineering blog. [linkedin.com/blog/engineering/feed/understanding-feed-dwell-time](https://engineering.linkedin.com/blog/engineering/feed/understanding-feed-dwell-time)
6. Tu Y, Lo C, Yuan Y, Chatterjee S (2019) Feedback shaping: a modeling approach to nurture content creation. *Proc. 25th ACM SIGKDD Conf. on Knowledge Discovery and Data Mining*. arxiv.org/abs/2106.11312
7. Agarwal D, Basu K, Ghosh S, *et al.* (2018) Online parameter selection for web-based ranking problems. *Proc. 24th ACM SIGKDD Conf. on Knowledge Discovery and Data Mining*, pp 23–32. doi.org/10.1145/3219819.3219847
8. Borisyuk F, Zhou M, Song Q, *et al.* (2024) LiRank: industrial large scale ranking models at LinkedIn. [arXiv:2402.06859](https://arxiv.org/abs/2402.06859). arxiv.org/abs/2402.06859
9. LinkedIn Engineering (2019) Leveraging dwell time to improve member experiences on the LinkedIn feed. Engineering blog. [linkedin.com/blog/engineering/feed/leveraging-dwell-time-to-improve-member-experiences](https://engineering.linkedin.com/blog/engineering/feed/leveraging-dwell-time-to-improve-member-experiences)
10. LinkedIn Engineering (2020) Engineering the next generation of LinkedIn’s feed. Engineering blog. [linkedin.com/blog/engineering/feed/engineering-the-next-generation-of-linkedins-feed](https://engineering.linkedin.com/blog/engineering/feed/engineering-the-next-generation-of-linkedins-feed)
11. Practitioner estimates of engagement weighting, dwell thresholds, and posting cadence are drawn from recurring independent analyses of observed post performance, of which the most widely circulated is the annual *Algorithm Insights* report series. These are uncontrolled observational studies of varying sample composition; where they disagree we report the range rather than a point estimate. All such figures are marked  in Table 2.
12. X Corp. (2023) The Algorithm: source code for the X recommendation algorithm. GitHub repository. Cited for the comparison of disclosure regimes in the Introduction and Conclusion. github.com/twitter/the-algorithm