

# Morgan Linton

Public profile snapshot

<https://www.linkedin.com/in/morganlinton>

Retrieved 2026-09-14T19:44:48.415Z · 7 public posts · Profile: Bright Data · Posts: Bright Data

## Profile presentation grade

100/100 · Grade A+ · 45% coverage · Partial grade. Rubric profile-2026.09.1.

A transparent presentation checklist, not a LinkedIn ranking, talent assessment or prediction of reach.

Grade thresholds: A+ 90, A 80, B 65, C 50, D 35, F below 35.

Score = earned points ÷ assessed points × 100. Coverage is assessed weight out of 100. Unavailable and masked data are excluded. Scores with different coverage are not directly comparable.

Truncated text earns only checks already demonstrated in the visible excerpt; the unseen remainder is not marked wrong. Text rules measure structure, not factual accuracy, writing quality or every language equally.

Photo and banner grades cover image availability, decoding and shape only. Face framing, lighting, expression, visual style and authenticity require a human review. No appearance, identity or personality judgments are made.

An education entry is evidence of profile context, not verified credentials. School prestige, age, degree level and background do not affect the score.

Profile picture: A+ · 100/100 · 15/15 points assessable

Profile image supplied: 5/5 points

The source supplied an image URL.

Action: Choose a current, recognizable photo. Preview the circular crop at comment size and leave space around your head.

Image decodes and fits its display shape: 10/10 points

Returned image: 200 × 200px. Dimensions and decoding were measured on the returned image, which may be a resized LinkedIn thumbnail. Original resolution, face framing, lighting and professional suitability were not assessed automatically.

Action: Use a square original, check sharpness at full size, and preview the circular crop. A small provider thumbnail does not prove your original upload is low resolution.

Cover & visual identity: A+ · 100/100 · 10/10 points assessable

Cover image supplied: 5/5 points

The source supplied an image URL.

Action: Use a simple cover that says who you help and what you do. Keep important text clear of the profile-photo overlay.

Image decodes and fits its display shape: 5/5 points

Returned image: 798 × 200px. Dimensions and decoding were measured on the returned image, which may be a resized LinkedIn thumbnail. Original resolution, face framing, lighting and professional suitability were not assessed automatically.

Action: Preview your cover on desktop and mobile. Start with a 1584 × 396 canvas and keep the message readable without tiny text.

Name & identity: A+ · 100/100 · 10/10 points assessable

Readable display name: 5/5 points

Source text: "Morgan Linton"

Action: Use the name people know you by, consistently across your profile and other professional channels. Single names and all writing systems are supported.

Name field stays focused on identity: 5/5 points

Source text: "Morgan Linton"

Action: Move URLs, contact details and sales calls to action into About or Contact info. Keep your actual name unchanged.

Headline & positioning: Not assessed · —/100 · 0/15 points assessable

Headline supplied: Unknown

Not returned in readable form by the source.

Action: State your role or specialty in the headline, using terms your intended reader understands.

Enough context to explain your work: Unknown

Not returned in readable form by the source.

Action: Use this structure: [role or specialty] · Helping [audience] achieve [outcome]. Replace every bracket with a truthful detail.

Headline is concise enough to scan: Unknown

Not returned in readable form by the source.

Action: Keep the headline focused on one main promise. Move your longer story into About.

About & descriptions: A+ · 100/100 · 5/20 points assessable

Readable introduction: 5/5 points

Truncated source excerpt: "a little about me, in Gonzo journalistic style...because that's just more fun..."

Action: Lead with who you help, the problem you solve and the work you do today.

An introduction with supporting detail: Unknown

Truncated source excerpt: "a little about me, in Gonzo journalistic style...because that's just more fun..."

Action: Add a specific example of your work, your approach, and the kinds of opportunities you welcome. Length is a structure check, not a quality guarantee.

Easy-to-scan paragraph structure: Unknown

Truncated source excerpt: "a little about me, in Gonzo journalistic style...because that's just more fun..."

Action: Break About into short paragraphs: current focus, relevant proof, then how to connect.

A clear next step: Unknown

Truncated source excerpt: "a little about me, in Gonzo journalistic style...because that's just more fun..."

Action: End with one relevant next step, such as who should message you and what to include. This automated check recognizes English phrases and links; review other languages manually.

Experience & proof: Not assessed · —/100 · 0/15 points assessable

Professional experience identified: Unknown

No readable entries were returned. We cannot conclude this section is empty.

Action: List relevant roles with an accurate title and organization. Connect them to the work you want to do next.

Experience includes useful descriptions: Unknown

Not returned in readable form by the source.

Action: For each relevant role, explain your responsibility, a concrete contribution, and an outcome you can substantiate. Do not invent numbers. Masked source descriptions need manual review.

Education & learning: Not assessed · —/100 · 0/10 points assessable

Education or learning entry identified: Unknown

No readable entries were returned. We cannot conclude this section is empty.

Action: Add relevant education or structured learning you want to share. Name the institution or course accurately; formal degrees are not required for a strong professional story.

Learning context is described: Unknown

Not returned in readable form by the source.

Action: Describe your subject, course or qualification and one relevant project or skill. Use the full qualification name where helpful. We assess clarity, never institution prestige or degree level.

Location & context: A+ · 100/100 · 5/5 points assessable

Readable location: 5/5 points

Source text: "Incline Village"

Action: Use an accurate city or region if you want it visible. Do not add a home address. Omitted provider data is not penalized.

## Executive summary

Morgan Linton's audit covers 7 public posts, including 7 with both reaction and comment counts. The measured median is 7 interactions. AI & building products is the most represented detected theme (4/7; themes can overlap). Start with a concrete follow-up to an existing subject, then test the opening and reply invitation across a consistent four-week sample.

## Measured performance

Followers: 5,424. Posts: 7. Fully measured: 7.

Engagement by current followers: 1.2%. Average reactions: 56.3. Median reactions: 3.

Average comments: 9. Comments per 100 reactions: 16.

Observed cadence: 8.4 posts/week. Sample window: 5 days.

## Strengths, watch points and context

STRENGTH: A usable personal baseline

7 of 7 posts have both public counts. Their median is 7 interactions. Compare future experiments with your own sample rather than an unrelated peer score.

Evidence: post 1, post 2, post 3, post 4, post 5, post 6, post 7. Full text and links are in Source posts.

STRENGTH: A concrete post to learn from

"Very cool waking up to see Shopify acquired Tailwind. If you don't know much about Tailwind, here's a quick TL;DR" has 389 visible interactions. That is 55.6× the sample median. Its topic and opening are candidates for a follow-up experiment; this does not identify why it performed.

Evidence: post 7. Full text and links are in Source posts.

CONTEXT: Conversation is observable; quality is not

2 of 7 posts contain a question. Measured posts generated 16 comments per 100 reactions. Counts can include the author's replies and do not measure lead quality or sentiment.

Evidence: post 4, post 7. Full text and links are in Source posts.

CONTEXT: Understand the gaps in this sample

The median gap between sampled dates is 0.9 days; the longest is 1.1 days (6 intervals). Missing posts can exaggerate gaps, so this is not a complete publishing calendar.

Evidence: post 7, post 6, post 5, post 4, post 3, post 2, post 1. Full text and links are in Source posts.

## Performance and publishing rhythm

Median visible interactions: 45 for the earliest 3 measured posts and 3 for the latest 3. The middle post is excluded to keep group sizes equal. This compares selected posts of different ages, not account growth.

Typical measured interactions: 7. Highest: 389.

Median gap 0.9 days; longest 1.1 days across 6 intervals.

2026-09-09T15:15:37.513Z: 365 reactions + 24 comments = 389 interactions.

2026-09-10T13:58:49.634Z: 2 reactions + 1 comments = 3 interactions.

2026-09-11T16:59:56.388Z: 13 reactions + 32 comments = 45 interactions.

2026-09-12T13:59:35.871Z: 3 reactions + 4 comments = 7 interactions.

2026-09-13T15:13:51.748Z: 6 reactions + 1 comments = 7 interactions.

2026-09-14T04:20:28.498Z: 2 reactions + 1 comments = 3 interactions.

2026-09-14T15:06:30.941Z: 3 reactions + 0 comments = 3 interactions.

UTC publishing days — descriptive, not recommended times:

Sun UTC: 1 posts / 1 measured; median 7, average 7 interactions.

Evidence: post 3. Full text and links are in Source posts.

Mon UTC: 2 posts / 2 measured; median 3, average 3 interactions.

Evidence: post 2, post 1. Full text and links are in Source posts.

Wed UTC: 1 posts / 1 measured; median 389, average 389 interactions.

Evidence: post 7. Full text and links are in Source posts.

Thu UTC: 1 posts / 1 measured; median 3, average 3 interactions.

Evidence: post 6. Full text and links are in Source posts.

Fri UTC: 1 posts / 1 measured; median 45, average 45 interactions.

Evidence: post 5. Full text and links are in Source posts.

Sat UTC: 1 posts / 1 measured; median 7, average 7 interactions.

Evidence: post 4. Full text and links are in Source posts.

Older posts have had more time to accumulate interactions. Different publication-day groups have different sample sizes and topics.

## Content pillars and format performance

Overlapping English keyword themes, not semantic certainty:

AI & building products: 4 posts / 4 measured; median 26, average 111 interactions.

Evidence: post 2, post 3, post 5, post 7. Full text and links are in Source posts.

Other / not classified by these rules: 3 posts / 3 measured; median 3, average 4.3 interactions.

Evidence: post 1, post 4, post 6. Full text and links are in Source posts.

Careers & leadership: 1 posts / 1 measured; median 389, average 389 interactions.

Evidence: post 7. Full text and links are in Source posts.

Source-reported formats; unknown is not text-only:

image: 5 posts / 5 measured; median 7, average 89.4 interactions.

Evidence: post 1, post 4, post 5, post 6, post 7. Full text and links are in Source posts.

text: 2 posts / 2 measured; median 5, average 5 interactions.

Evidence: post 2, post 3. Full text and links are in Source posts.

Small or unequal groups cannot establish a winning format.

## Hooks and writing structure

Play more video games.

Hook: Statement; 4 words; 1 paragraphs; final CTA not detected; question not detected.

Hashtags: None detected. Mentions: None detected.

Very interesting week ahead. There will be two main topics discussed:

Hook: Statement; 52 words; 5 paragraphs; final CTA not detected; question not detected.

Hashtags: None detected. Mentions: None detected.

We are at a time in history where AI safety could possibly be most important way to make an impact in the world.

Hook: Personal opening; 172 words; 7 paragraphs; final CTA detected; question not detected.

Hashtags: None detected. Mentions: None detected.

So I had a post go super viral on LinkedIn, but I'm not a marketer, I'm an engineer.

Hook: Statement; 55 words; 3 paragraphs; final CTA not detected; question detected.

Hashtags: None detected. Mentions: None detected.

Totally shocking news, all these AI labs like Kimi, Qwen, and DeepSeek, that we thought had found a way to provide open weight models that rivaled OpenAI and Anthropic...

Hook: Statement; 81 words; 3 paragraphs; final CTA not detected; question not detected.

Hashtags: None detected. Mentions: None detected.

Okay, it took a week to get this done the right way, but it's finally all complete, my comparison of Astra and Fable 5.1 with VulcanBench [U+1F596]

Hook: Number-led; 277 words; 7 paragraphs; final CTA not detected; question not detected.

Hashtags: None detected. Mentions: None detected.

Very cool waking up to see Shopify acquired Tailwind. If you don't know much about Tailwind, here's a quick TL;DR

Hook: Statement; 202 words; 9 paragraphs; final CTA not detected; question detected.

Hashtags: None detected. Mentions: None detected.

Hook uses the first nonempty line. Paragraphs use blank-line breaks. Final CTA uses the last content line, ignoring tag-only lines, with question or action-phrase detection. Text rules can miss meaning and linked mentions.

## Hashtag patterns

Up to 20 most-used tags. Each post is counted once per tag. Tags travel with topics and formats; their isolated effect cannot be measured.

## Prioritized action plan

### 1. Develop a specific follow-up

Observation: The sample includes “Very cool waking up to see Shopify acquired Tailwind. If you don't know much about Tailwind, here's ”. Its visible interaction total is 389.

Action: Build a follow-up to “Very cool waking up to see Shopify acquired Tailwind. If you don't know much about Tailwind, here's ”. Add a concrete example, a decision you made and one lesson. Keep the format similar so topic development is the main experiment.

Potential impact: Test repeatable interest in an existing subject. Effort: Medium.

Metric to watch: Reactions and comments after the same seven-day window; compare with the baseline median

Evidence: post 7. Full text and links are in Source posts.

### 2. Test a sharper opening

Observation: Median text length is 81 words with 5 paragraphs. 5 of 7 openings are classified as statements.

Action: For the next post related to “Very cool waking up to see Shopify acquired Tailwind. If you don't know much about Tailwind, here's ”, open with a specific problem and a concrete outcome you can substantiate. Move context into the second paragraph. Avoid invented results or curiosity gaps that the post does not resolve.

Potential impact: Improve clarity; any engagement effect remains to be tested. Effort: Low.

Metric to watch: Median interactions across at least three comparable posts per opening approach

Evidence: post 1, post 2, post 3. Full text and links are in Source posts.

### 3. Make the reply invitation concrete

Observation: 1 of 7 final content lines contain a detected question or action phrase; 2 posts contain a question anywhere. These are text rules, not semantic judgments.

Action: End one upcoming post with a focused trade-off question connected to ai & building products. Ask readers to share a decision and why. Read replies before deciding whether the discussion was useful.

Potential impact: Test whether a clearer invitation supports useful discussion. Effort: Low.

Metric to watch: Public comments per post and manual review of reply relevance

Evidence: post 3. Full text and links are in Source posts.

#### 4. Choose a sustainable publishing experiment

Observation: 7 dated posts span 5 days. Observed cadence is 8.4 posts per week.

Action: Choose two slots per week for the four-week experiment below, or reduce the schedule if your capacity is lower. Keep the slots consistent and log every published post. Two slots are a planning choice, not a LinkedIn benchmark.

Potential impact: Create comparable records and a manageable learning loop. Effort: Medium.

Metric to watch: Planned versus published posts, time spent, and median interactions at seven days

Evidence: post 7, post 6, post 5, post 4, post 3, post 2, post 1. Full text and links are in Source posts.

#### 5. Use tags and mentions with a purpose

Observation: 0 of 7 posts contain hashtags and 0 contain visible @mentions. LinkedIn may return linked mentions as plain text, so detection can undercount.

Action: Use relevant topic labels only when they clarify the subject. Mention people only when they contributed or are directly relevant. Test a consistent tagging approach on comparable posts; do not attribute performance changes to tags alone.

Potential impact: Keep context and attribution useful. Effort: Low.

Metric to watch: Median interactions for tagged/untagged samples, with sample sizes and topic differences

#### 6. Test presentation without changing the subject

Observation: image: 5 posts (5 measured); text: 2 posts (2 measured).

Action: Turn one idea related to "Very cool waking up to see Shopify acquired Tailwind. If you don't know much about Tailwind, here's " into a concise checklist or document post. Keep a similar topic and ask the same type of question. Label this a new format experiment, not a proven winning format.

Potential impact: Learn whether a different presentation helps explain the same idea. Effort: Medium.

Metric to watch: Public interactions at the same post age; at least three posts in each format before comparison

Evidence: post 1, post 2, post 3, post 4, post 5, post 6, post 7. Full text and links are in Source posts.

## Basic content ideas

### 1. The next chapter

Follow up on "Very cool waking up to see Shopify acquired Tailwind. If you don't know much about Tailwind, here's ": what changed since publishing, what stayed difficult, and the next step. Only include outcomes you can verify.

Evidence: post 7. Full text and links are in Source posts.

### 2. Show the process

Break the work behind "Very cool waking up to see Shopify acquired Tailwind. If you don't know much about Tailwind, here's " into three decisions. Add an example or screenshot you have permission to share.

Evidence: post 7. Full text and links are in Source posts.

### 3. A useful trade-off

Describe two approaches to ai & building products, explain your constraint, and ask which constraint readers face.

Evidence: post 2, post 3, post 5, post 7. Full text and links are in Source posts.

## Basic writing outline

Open with one specific point. Show your reasoning with a supported example. End with a focused question.

Peer comparison, private AI strategy, extensive schedules and downloadable writing skills are coming with Profilyst Pro. They are not included in this public report.

## Source posts

### POST 1

[https://www.linkedin.com/posts/morganlinton\\_play-more-video-games-activity-7505280828722294784-dZZd](https://www.linkedin.com/posts/morganlinton_play-more-video-games-activity-7505280828722294784-dZZd)

Published: 2026-09-14T15:06:30.941Z. Reactions: 3. Comments: 0. Format: image.

Play more video games.

### POST 2

[https://www.linkedin.com/posts/morganlinton\\_very-interesting-week-ahead-there-will-be-activity-7505118247252283392-hhok](https://www.linkedin.com/posts/morganlinton_very-interesting-week-ahead-there-will-be-activity-7505118247252283392-hhok)

Published: 2026-09-14T04:20:28.498Z. Reactions: 2. Comments: 1. Format: text.

Very interesting week ahead. There will be two main topics discussed:

1. Pacing AI model releases

2. A bunch of new AI model releases

Welcome to September 2026.

Oh, and AI stocks drop like a rock first thing Monday morning. I don't make the rules, that's just going to happen.

Buckle up.

POST 3

[https://www.linkedin.com/posts/morganlinton\\_we-are-at-a-time-in-history-where-ai-safety-activity-7504920289735118849-U0PJ](https://www.linkedin.com/posts/morganlinton_we-are-at-a-time-in-history-where-ai-safety-activity-7504920289735118849-U0PJ)

Published: 2026-09-13T15:13:51.748Z. Reactions: 6. Comments: 1. Format: text.

We are at a time in history where AI safety could possibly be most important way to make an impact in the world.

As one human, I'm trying to increase the impact I can make. And I think there's an opportunity here, that I just can't stop thinking about.

I am starting to put the pieces together on an eval suite for VulcanBench, focused on AI safety.

It's a different approach from what companies like METR are taking. And I'm not diminishing what they are doing in any way, but I am saying there is room for other approaches.

I want to look at AI safety, as companies use AI today, in the actual harnesses they use, with the actual things, their teams are using AI for every single day.

It is very meaningful for me to be able to make these kinds of free, open source, contributions. This is what open source is all about, doing things to make an impact above all else.

More details in the first comment below.

POST 4

[https://www.linkedin.com/posts/morganlinton\\_so-i-had-a-post-go-super-viral-on-linkedin-activity-7504539212566917120-4chX](https://www.linkedin.com/posts/morganlinton_so-i-had-a-post-go-super-viral-on-linkedin-activity-7504539212566917120-4chX)

Published: 2026-09-12T13:59:35.871Z. Reactions: 3. Comments: 4. Format: image.

So I had a post go super viral on LinkedIn, but I'm not a marketer, I'm an engineer.

Anyone following me know about how to analyze this so I can know what I did right?

Right now I'm like, let's get more out like this!! But then I'm also wondering, what does that even mean [U+1F633]

POST 5

[https://www.linkedin.com/posts/morganlinton\\_totally-shocking-news-all-these-ai-labs-activity-750422209243066368-IEHN](https://www.linkedin.com/posts/morganlinton_totally-shocking-news-all-these-ai-labs-activity-750422209243066368-IEHN)

Published: 2026-09-11T16:59:56.388Z. Reactions: 13. Comments: 32. Format: image.

Totally shocking news, all these AI labs like Kimi, Qwen, and DeepSeek, that we thought had found a way to

provide open weight models that rivaled OpenAI and Anthropic...

It looks like they were actually cheating, and just serving Claude models in the backend the whole time.

I still need to do a deep dive here but if this does turn out to be true, it kinda changes the entire narrative on local ai and the pace it was actually moving.

POST 6

[https://www.linkedin.com/posts/morganlinton\\_okay-it-took-a-week-to-get-this-done-the-activity-7503814242903461888-5Ct9](https://www.linkedin.com/posts/morganlinton_okay-it-took-a-week-to-get-this-done-the-activity-7503814242903461888-5Ct9)

Published: 2026-09-10T13:58:49.634Z. Reactions: 2. Comments: 1. Format: image.

Okay, it took a week to get this done the right way, but it's finally all complete, my comparison of Astra and Fable 5.1 with VulcanBench [U+1F596]

A few things took longer here, the primary one being some updates to my benchmarking score to layer in code quality. This added 3-4 days of time, but for the right reasons.

This new class of models requires a new class of benchmarks. I don't think we can just look at things like accuracy any more, we also have to look at code quality/maintainability factors.

Now with VulcanBench-SWE v4, 33% of the score is code quality/maintainability. For me, and many other engineering leaders, seeing a model get a 98% on a benchmark doesn't really give us much signal.

I created VulcanBench to help make decisions around model and effort level, and this means not just building evals that represent the kind of work teams give to these models, but the kind of output we expect from these models when building scalable system and working in large codebases.

While I would normally share more about my thoughts, I'll let you come to your own conclusions about Astra and Fable 5.1. Both are excellent models, OpenAI and Anthropic have really created a new class of models here, now it's for us to decide if we need this horsepower for daily tasks, or just the hard stuff, and to be realistic about the quality of the output.

Model card below, and if you want to do a deep dive, you can find more on the VulcanBench site and Github, both shared in the first comment below.

Live long and benchmark [U+1F596]

POST 7

[https://www.linkedin.com/posts/morganlinton\\_very-cool-waking-up-to-see-shopify-acquired-activity-7503471181883314176-Qwxk](https://www.linkedin.com/posts/morganlinton_very-cool-waking-up-to-see-shopify-acquired-activity-7503471181883314176-Qwxk)

Published: 2026-09-09T15:15:37.513Z. Reactions: 365. Comments: 24. Format: image.

Very cool waking up to see Shopify acquired Tailwind. If you don't know much about Tailwind, here's a quick TL;DR

Tailwind is one of the most polished and well-built CSS frameworks, i.e. makes sites look really nice.

They're also kinda famous for having an insanely brilliant team.

So why did Shopify buy them?

I think it's for a few reasons, and many of these are related to the fact that Tobias Lütke really understand technology, and where things are going with agentic commerce.

Shopify has been really jamming with headless commerce w/Hydrogen, and Tailwind is the main styling tool here.

Hydrogen is a React-based framework, and React is all about easy, plug-and-play reusable components. It's super easy for developers to pair Tailwind UI libraries w/Hydrogen so they can essentially copy and paste modular layouts.

Couple this with the fact that agentic coding workflows though Claude, OpenAI, Grok, etc. all bias towards Tailwind, and it's honestly one of the most well-timed/best-fit acquisitions I've seen.

Once again, Tobi is a genius, and yeah, he now has another team of geniuses joining Shopify. I would not like to be a Shopify competitor right now, they are getting really far ahead.

## Saved audit comparison

This export represents one saved snapshot. No earlier comparable snapshot was selected.

## Methodology and limits

Read 7 of up to 8 requested posts. Unreturned or unverifiable posts are not included in the analysis.

The source returned a shortened About section. Suggestions use only the visible excerpt.

Some profile sections were not returned by the source. They remain unknown, not confirmed missing from LinkedIn.

Selected URLs may omit recent or low-performing posts. This is a sample, not a complete feed.

Reaction and comment totals are snapshots. Older posts had more time to accumulate interactions.

Themes can overlap and use English keyword rules. Hooks, CTAs and mentions are text heuristics and may miss context or other languages.

Timing is UTC publishing history, not audience activity or an optimal posting time.

Recommendations and potential impact are hypotheses. No impressions, revenue, private demographics, peer percentile or causal attribution are available.

Engagement = average visible reactions plus comments / current followers × 100, requiring at least three measured posts. Missing values remain unavailable. Cadence uses the intervals between sampled dates. No uncalibrated score or private metric is inferred.

Independent public-data analysis. Not affiliated with LinkedIn or the profile owner. <https://profilyst.com/methodology>

Unsupported glyphs are preserved as Unicode code points [U+...]. Full original text is available in the online report.