

Temporal Changes in Hacker News Hiring Posts Around COVID-19 and the Post-ChatGPT Period

Ruby, Rails, Elixir, and comparator ecosystems, 2019–2026

Marian Posăceanu
Independent Researcher

19 July 2026

Working paper; not peer reviewed; author validation and final sign-off pending

Abstract. Monthly Hacker News *Ask HN: Who is hiring?* threads form a recurring observational record of employer communication in a technically oriented labor market. This paper examines posting volume, remote-work terminology, and mentions of Ruby, Ruby on Rails, and Elixir around COVID-19 and the public release of ChatGPT. Formal analysis uses a consistent aggregate lexical panel of 40,087 top-level posts across 65 months, January 2019–May 2024. A separately measured post-level extension supports aggregate overlap checks and descriptive coverage through July 2026. Count outcomes are modeled with a negative-binomial segmented regression; term shares use grouped-binomial models. All principal models include linear trend, seasonal indicators, COVID pulse and adaptation terms, January 2023 level and slope terms, HAC uncertainty, and false-discovery-rate correction. Comparator indices cover 11 other languages and 10 other frameworks. The acute COVID interval is associated with 24.6% lower posting volume and substantially higher odds of remote terminology. Ruby and Rails odds decline during the acute interval, but comparator-adjusted COVID gaps are not statistically clear. January 2023 is associated with lower target-stack odds and a 35.9% lower Ruby/Elixir language index relative to comparator languages. However, Ruby, Rails, and both comparator gaps were already declining during 2021–2022; nearby candidate break dates often produce similarly negative estimates; and within the 17-month post-launch window neither broad-AI nor LLM-term intensity is significantly correlated with target shares after detrending. The evidence supports a durable COVID-era remote-work transition and a later relative decline in the target ecosystems, but it does not identify direct displacement by LLMs. Because post-level human precision and recall validation remains incomplete, the results should be treated as exploratory.

Keywords: software hiring; Hacker News; Ruby; Ruby on Rails; Elixir; COVID-19; remote work; artificial intelligence; large language models; interrupted time series.

1 Introduction

Online vacancy text can reveal changes in employer demand before those changes appear in conventional occupational statistics, but it is also a selective and noisy measurement channel. Job advertisements combine recruitment, branding, skill signalling, and administrative boilerplate; one advertisement can represent several roles, while repeated advertisements need not imply new vacancies. Prior research nevertheless shows that large collections of job postings can measure changes in skill requirements and employer composition when the unit and selection mechanism are made explicit [1–3].

The monthly *Ask HN: Who is hiring?* thread offers a smaller but unusually regular corpus. Employers post free-text comments to a recurring thread, commonly naming roles, technologies, locations, and work arrangements. HNTrends has already

documented broad descriptive movements in the archive, including the 2020 posting contraction, the remote-work shift, the 2023 hiring downturn, and the rise of AI terminology [4]. The contribution here is narrower: a formal, reproducible examination of Ruby, Rails, and Elixir; explicit comparison with other language and framework ecosystems; and a separation between descriptive timing and causal attribution.

Two events motivate the analysis. WHO characterized COVID-19 as a pandemic on 11 March 2020 [5]. The April 2020 hiring thread is the first complete monthly thread after that announcement. OpenAI introduced ChatGPT publicly on 30 November 2022 [6]; January 2023 is the first complete monthly thread after launch. Neither date is treated as a randomized intervention. The second date, in particular, coincides with a technology-sector hiring contraction, monetary tightening, financing changes, and pre-existing stack trends.

The paper addresses four questions:

RQ1. How did total HN hiring-post volume and remote-work terminology change around the initial COVID disruption and subsequent adaptation period?

RQ2. Did Ruby, Rails, and Elixir change differently from comparator language and framework ecosystems?

RQ3. Are January 2023 target-stack changes robust to pretrend diagnostics, alternative break dates, and LLM-specific rather than broad-AI measures?

RQ4. Does a separately measured extension through 2026 support the direction of the formal-panel results?

The central finding is asymmetrical. The COVID-to-remote transition is large, abrupt, and persistent. The post-2022 decline in target-stack visibility is also measurable, including relative to comparator languages, but significant pre-event declines and break-date non-uniqueness prevent a direct claim that ChatGPT or LLM adoption caused it.

2 Related Work

2.1 Vacancy text as labor-market evidence

Vacancy postings provide detailed, employer-written descriptions of desired skills, but their interpretation depends on platform selection and posting behavior. Deming and Kahn use professional job ads to study within-occupation variation in cognitive and social skill requirements [1]. Hershbein and Kahn show that recessions can coincide with persistent upgrading of vacancy requirements, illustrating why a market downturn can alter both posting volume and composition [2]. Acemoglu and colleagues use establishment-level online vacancies to examine the growth of AI-related demand and changes in non-AI hir-

ing at AI-exposed establishments [3]. These studies motivate treating posting text as an outcome in its own right rather than equating mentions with completed hires.

2.2 COVID-19, work from home, and software work

COVID-19 generated a rapid reallocation toward work that could be performed remotely. Dingel and Neiman classify occupational work-from-home feasibility [7], while Barrero, Bloom, and Davis document the broader persistence of work from home after the initial shock [8]. Software-engineering studies report heterogeneous developer experiences, collaboration changes, and well-being effects during pandemic remote work [9, 10]. The present study does not measure productivity or worker welfare; it measures employer use of remote terminology in a recurrent hiring channel.

2.3 Generative AI, productivity, and employment

Experimental and field evidence shows that generative-AI assistance can improve performance in some bounded writing, support, and programming tasks [11–13]. Such productivity effects do not mechanically imply fewer vacancies, and effects can vary by task, worker experience, and organizational adoption. The labor-demand question is therefore distinct from the task-productivity question. This paper contributes only a descriptive ecosystem-level time series. It lacks firm-level adoption, exposure, or an untreated counterfactual, so it cannot estimate employment displacement.

2.4 Interrupted time series and multiplicity

Segmented interrupted-time-series models are useful for describing level and slope changes around defined events, but they require careful treatment of autocorrelation, seasonality, pre-existing trends, concurrent events, and functional form [14]. The analysis uses heteroskedasticity-and-autocorrelation-consistent covariance estimates [15] and controls the false discovery rate across prespecified event tests using the Benjamini–Hochberg procedure [16]. Pretrend and date-sensitivity analyses are treated as substantive diagnostics rather than appendix formalities.

3 Data and Measurement

3.1 Formal aggregate panel

The formal sample is a frozen HNTrends lexical panel covering January 2019–May 2024: 65 monthly threads and 40,087 top-level comments. Each monthly observation contains the total top-level comment count and term-specific counts and percentages. The panel does not include post-level text or employer identifiers. HNTrends is an upstream descriptive source. Because redistribution rights for its source snapshot were not established, the public package supplies the selected analytical panel and a checksum/provenance record rather than the upstream snapshot itself [4].

The primary target terms are Ruby, Rails, and Elixir. Remote and broad AI are modeled separately. The broad-AI series is the upstream AI share; it predates generative AI and should not be read as an LLM-adoption measure. A distinct LLM-term intensity sums the percentages for GPT, GPT-3, GPT-4, ChatGPT, and OpenAI. Because a post can mention several terms, additive intensities can exceed the share of unique posts mentioning any one concept.

3.2 Comparator construction

The language target set is Ruby and Elixir. Its comparator set is Python, JavaScript, TypeScript, Go, Java, Rust, PHP, C#, Kotlin, Scala, and Clojure. The framework target set is Rails and Phoenix. Its comparators are Django, Spring, Laravel, React, Node.js, Express, ASP.NET, Vue, AngularJS, and Svelte.

For each term j , monthly share p_{jt} is normalized by its January 2019–December 2022 mean. A 0.05-percentage-point pseudocount prevents zero-valued series from dominating log transformations. The group index is the geometric mean of normalized term series:

$$I_{Gt} = 100 \exp \left[|G|^{-1} \sum_{j \in G} \log \frac{p_{jt} + 0.05}{\bar{p}_{j,2019:2022} + 0.05} \right]. \quad (1)$$

The comparator gap is $g_t = \log(I_{target,t}/I_{comp,t})$. This construction gives each named technology equal multiplicative weight; it is a sensitivity device, not a market-share portfolio.

3.3 Recent post-level extension

The later extension begins with committed monthly records in the open-source `hacker-job` project, which structures HN hiring comments using an LLM-assisted pipeline [17]. Additional matching and false-positive rules identify Ruby, Rails/Ruby on Rails/RoR, and Elixir. The extension spans July 2023–July 2026 and contains 695 retained matches, of which 692 are marked core matches. July 2026 is partial.

The public release removes original comment text, usernames, contact-email fields, external application URLs, and target-context excerpts. It retains HN comment IDs, canonical source URLs, and derived fields. This extension is used for aggregate agreement and descriptive recent-period summaries only. It is not appended to the formal regression panel.

3.4 Measurement validity

The two pipelines overlap for 11 months, July 2023–May 2024. Monthly count correlations are approximately 0.96 for each target term, with mean absolute count errors below one post per month (Table 2). This supports directional continuity, not post-level classification validity. A 225-row source-linked predicted-positive sample is distributed as a blinded annotation packet plus a separate key containing predictions, stratum sizes, inclusion probabilities, and weights. A separate helper can draw an all-post probability sample for recall from comments still retrievable through the official HN API; it treats the story’s direct child identifiers as top-level comments, consistent with the API data model [18]. The scorer refuses incomplete labels and does not report recall from a positive-only sample. Because the upstream aggregate panel does not expose its original per-post decisions, the helper validates a transparent reconstruction rather than reproducing the upstream classifier exactly. All human labels are intentionally blank in this release.

4 Empirical Strategy

4.1 Event definitions

The acute COVID indicator equals one for April–June 2020. A post-acute adaptation level and slope begin in July 2020. The January 2023 indicator marks the first full monthly thread after ChatGPT’s public launch; a corresponding post-event

Table 1. Analytical samples, uses, and unresolved validation tasks.

Sample	Period	Size	Analytical use	Principal limitation
HN Trends lexical panel	Jan. 2019–May 2024	65 months; 40,087 posts	Formal time-series models and comparator indices	Aggregate dictionary counts; no post-level confusion matrix in the frozen snapshot
Cleaned post-level extension	Jul. 2023–Jul. 2026	695 retained matches; 692 core matches	Overlap agreement and descriptive recent-period extension	Different extraction pipeline; July 2026 is partial; precision audit awaits author annotation
Public validation packet	Jul. 2023–Jul. 2026	225-row blinded stratified sample	Weighted precision audit using source HN URLs and a separate scoring key	Cannot estimate recall without a probability sample from all monthly posts

Table 2. Aggregate cross-source overlap, Jul. 2023–May 2024.

Term	Pearson r	Mean bias	MAE
Ruby	0.96	+0.55	0.73
Rails	0.96	+0.18	0.73
Elixir	0.96	+0.45	0.45

Notes: Monthly count agreement does not establish post-level precision or recall. The latter remains an explicit author-validation task.

slope captures subsequent change. These are time markers, not treatment assignments.

4.2 Count and share models

For total monthly posts N_t , the primary model is negative binomial with a log link. For each term count K_t out of N_t posts, the model is grouped binomial with a logit link. Both use the same design vector:

$$\eta_t = \alpha + \beta_1 t + \beta_2 \text{Acute}_t + \beta_3 \text{AdaptLevel}_t + \beta_4 \text{AdaptSlope}_t + \beta_5 \text{Jan2023}_t + \beta_6 \text{PostJan}_t + \sum_{m=2}^{12} \delta_m \text{Month}_{m,t}.$$

For counts, $\eta_t = \log E[N_t]$; for shares, $\eta_t = \text{logit } E[K_t/N_t]$. The negative-binomial dispersion parameter is estimated as 0.0152; Pearson χ^2/df is 1.36. Covariance estimates are HAC with six monthly lags. Reported count effects are incidence-rate changes; share effects are odds changes, not percentage-point changes.

The 24 prespecified primary tests comprise acute-COVID, post-acute-level, January-2023-level, and post-January-slope terms for six outcomes: total posts, Ruby, Rails, Elixir, remote, and broad AI. Benjamini–Hochberg q values control false discovery rate within this family. Comparator-gap event terms form a separate correction family.

4.3 Diagnostics and sensitivity analyses

Four analyses assess whether January 2023 should be interpreted as a distinctive break.

First, January 2021–December 2022 pre-period slopes are estimated with HAC(3). Second, the assumed break date is moved monthly from January 2021 through November 2023 while retaining the remaining design. The fraction of dates producing an estimate at least as negative as January 2023 is a descriptive rank, not a randomization p value. Third, 12-month pre/post windows are compared using a circular moving-block bootstrap with three-month blocks and 8,000 repetitions. Fourth, broad-AI and LLM-term intensity are correlated with target shares during January 2023–May 2024, including correlations after separately residualizing each series on linear time.

These checks do not create an untreated counterfactual. Their purpose is to detect overinterpretation of one chosen date or one model form.

5 Results

5.1 Posting volume and remote terminology

Figure 1 shows the sharp 2020 interruption, subsequent recovery, and lower 2023–2024 posting level. The negative-binomial model estimates 24.6% lower incidence during April–June 2020 (95% CI: 32.4% to 16.0% lower; $q < 0.001$). The post-acute level is 45.2% higher than the counterfactual continuation of the initial specification (95% CI: 11.2% to 89.8%; $q = 0.012$), consistent with recovery rather than a permanent 2020 collapse.

Remote terminology changes more strongly. The odds of a post mentioning remote work are 99.9% higher in the acute period and 249.3% higher at the post-acute adaptation level, both $q < 0.001$. The simpler 12-month comparison moves from 29.90% in April 2019–March 2020 to 62.75% in April 2020–March 2021, a 32.85-point increase (block-bootstrap 95% CI: 26.35 to 38.73 points). This is the clearest result in the study because it is large, persistent, and visible without dependence on a single fitted discontinuity.

The January 2023 model term is associated with 38.0% lower post incidence ($q < 0.001$). Remote odds are 51.1% lower at that level and continue declining by 4.7% per month in odds through May 2024, while remaining well above the pre-COVID baseline.

5.2 Target ecosystems around COVID-19

During the acute COVID interval, Ruby odds are 20.3% lower ($q = 0.031$) and Rails odds are 26.9% lower ($q < 0.001$). Elixir’s estimate is effectively zero with a wide interval. None of the target shares has a statistically clear post-acute adaptation-level change after multiplicity correction.

The comparator analysis further weakens a stack-specific COVID interpretation. The acute language target/comparator gap changes by -9.5% (95% CI: -29.6% to $+16.4\%$; $q = 0.584$), and the framework gap by -18.0% (95% CI: -40.9% to $+13.8\%$; $q = 0.471$). Ruby and Rails therefore dipped in absolute posting odds during the acute disruption, but the data do not show a clear persistent or uniquely target-stack COVID penalty relative to the wider comparison sets.

5.3 Broad AI and LLM-specific terminology

Broad AI language already appears throughout the pre-ChatGPT period, while GPT/ChatGPT/OpenAI terminology rises from a near-zero baseline after late 2022 (Figure 3). The January 2023 broad-AI level term is positive but not statistically clear

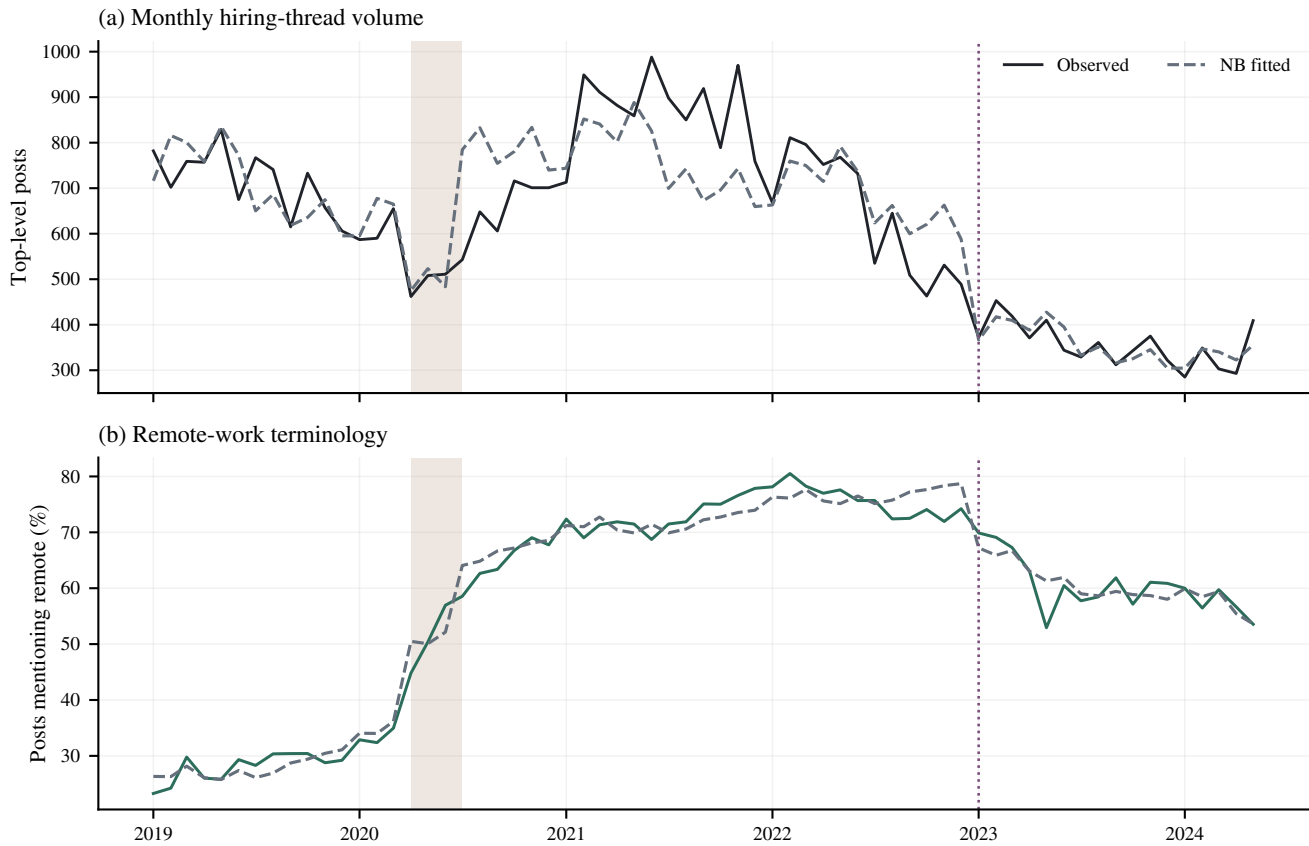


Figure 1. Monthly hiring-thread volume and remote-work terminology in the formal panel. Solid lines are observed values and dashed lines are fitted values. The shaded band is April–June 2020; the dotted line is January 2023. Fitted count values use the negative-binomial model, and fitted remote values use the grouped-binomial model.

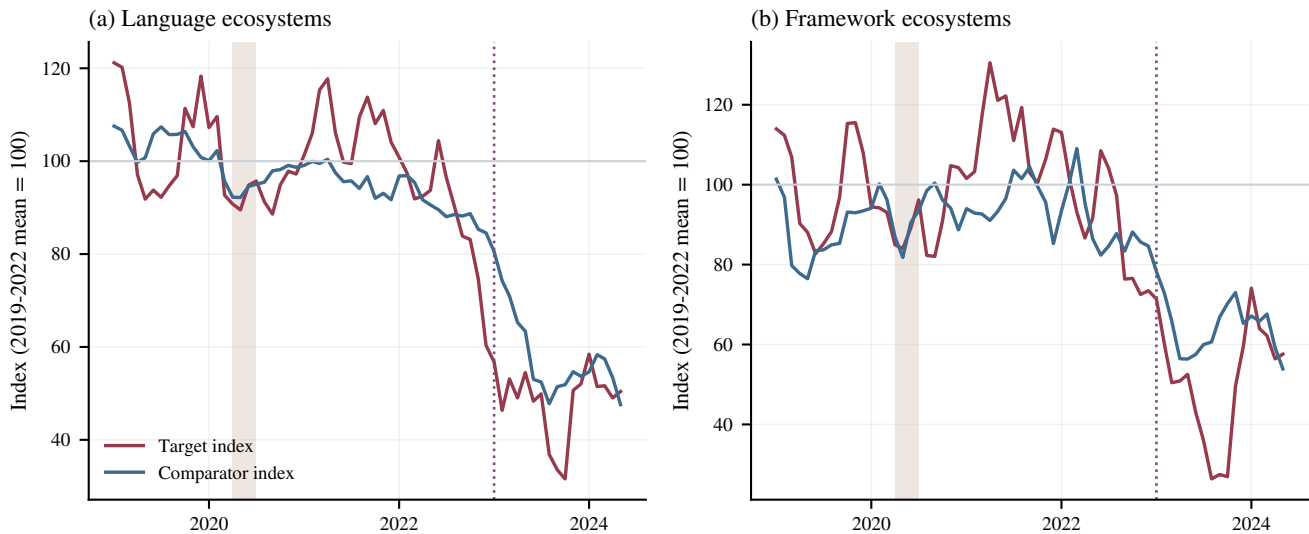


Figure 2. Normalized target and comparator ecosystem indices. Each named technology is normalized to its 2019–2022 mean, then aggregated by geometric mean. The shaded band is April–June 2020; the dotted line is January 2023. Indices describe relative visibility within this HN corpus, not labor-market employment shares.

after trend and seasonality (odds change +19.7%; $q = 0.185$). The subsequent broad-AI slope is strong: odds rise 5.66% per month (95% CI: 4.62% to 6.71%; $q < 0.001$).

This distinction matters. The broad-AI share combines older AI language and is unsuitable as a direct proxy for LLM adoption. The additive LLM-term series is more event-specific, but it is still a text measure rather than firm use of generative AI.

5.4 January 2023 target-stack estimates

At the January 2023 marker, Ruby odds are 31.6% lower, Rails odds 34.2% lower, and Elixir odds 61.7% lower; all three have $q < 0.001$. Figure 4 places these estimates beside the count, remote, and AI results. The model does not estimate statistically clear post-January target slopes, so the fitted effect is primarily a level difference within the chosen specification.

Table 3. Primary segmented models. Effects are incidence-rate changes for total posts and odds changes for term shares.

Outcome	Term	Effect (%)	95% CI	<i>p</i>	FDR <i>q</i>
Total hiring posts	Acute COVID	-24.6	[-32.4, -16.0]	< 0.001	< 0.001
Total hiring posts	Post-acute level	+45.2	[+11.2, +89.8]	0.006	0.012
Total hiring posts	Jan. 2023 level	-38.0	[-52.3, -19.5]	< 0.001	< 0.001
Total hiring posts	Post-Jan. 2023 slope	-0.6	[-2.1, +0.9]	0.450	0.515
Ruby share	Acute COVID	-20.3	[-33.9, -3.8]	0.018	0.031
Ruby share	Post-acute level	+2.5	[-14.0, +22.2]	0.779	0.813
Ruby share	Jan. 2023 level	-31.6	[-43.0, -17.9]	< 0.001	< 0.001
Ruby share	Post-Jan. 2023 slope	+0.4	[-1.1, +2.0]	0.587	0.641
Rails share	Acute COVID	-26.9	[-37.0, -15.3]	< 0.001	< 0.001
Rails share	Post-acute level	+9.4	[-5.7, +26.9]	0.236	0.316
Rails share	Jan. 2023 level	-34.2	[-45.5, -20.7]	< 0.001	< 0.001
Rails share	Post-Jan. 2023 slope	-1.0	[-2.6, +0.7]	0.237	0.316
Elixir share	Acute COVID	-0.1	[-23.3, +30.0]	0.993	0.993
Elixir share	Post-acute level	+10.6	[-12.3, +39.5]	0.395	0.474
Elixir share	Jan. 2023 level	-61.7	[-73.3, -45.2]	< 0.001	< 0.001
Elixir share	Post-Jan. 2023 slope	+1.5	[-1.5, +4.5]	0.331	0.419
Remote share	Acute COVID	+99.9	[+70.1, +134.9]	< 0.001	< 0.001
Remote share	Post-acute level	+249.3	[+189.4, +321.5]	< 0.001	< 0.001
Remote share	Jan. 2023 level	-51.1	[-63.9, -33.8]	< 0.001	< 0.001
Remote share	Post-Jan. 2023 slope	-4.7	[-5.8, -3.6]	< 0.001	< 0.001
AI share	Acute COVID	+23.9	[+9.1, +40.6]	< 0.001	0.002
AI share	Post-acute level	-11.3	[-25.6, +5.7]	0.180	0.269
AI share	Jan. 2023 level	+19.7	[-4.3, +49.7]	0.116	0.185
AI share	Post-Jan. 2023 slope	+5.7	[+4.6, +6.7]	< 0.001	< 0.001

Notes: Grouped-binomial models are used for shares; a negative-binomial model is used for total posts. All models include a linear trend, COVID acute pulse, post-acute level and slope, January 2023 level and slope, and month-of-year indicators. Covariance is HAC(6). FDR correction is Benjamini–Hochberg across the 24 primary event tests. Odds changes are not percentage-point changes.

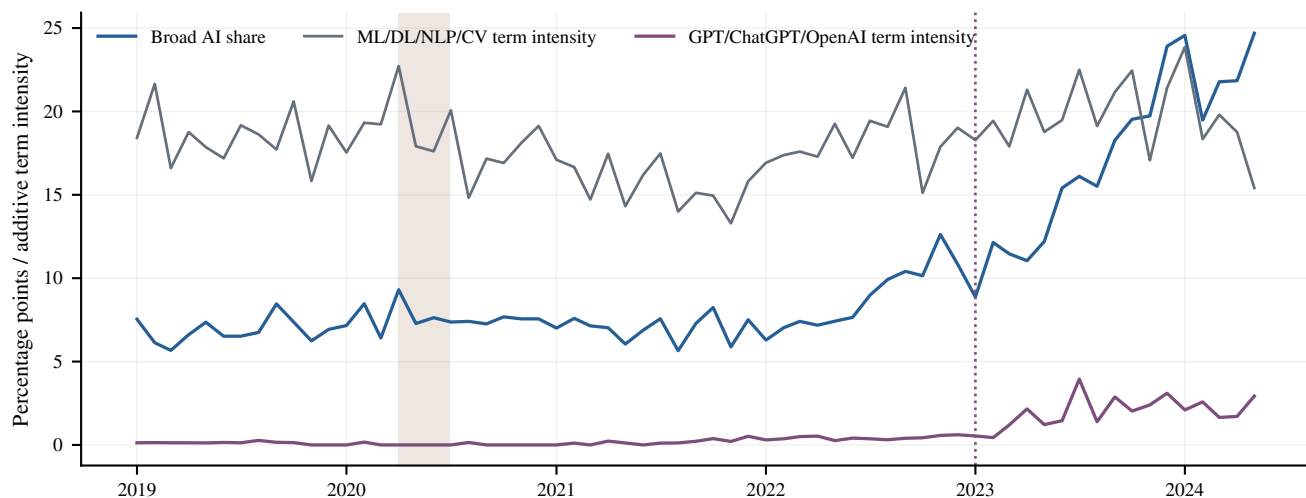


Figure 3. AI and LLM terminology in the formal panel. Broad AI is the upstream AI share. Traditional ML intensity sums machine learning, deep learning, NLP, and computer-vision percentages. LLM-term intensity sums GPT, GPT-3, GPT-4, ChatGPT, and OpenAI percentages and is therefore not a unique-post share.

Relative to comparator ecosystems, the language target index (Ruby/Elixir) is 35.9% lower at January 2023 (95% CI: 45.8% to 24.2% lower; $q < 0.001$), followed by a 2.82% monthly relative rebound. The framework target index (Rails/Phoenix) is 38.0% lower, but its interval reaches a small positive value and the FDR-adjusted result is not statistically clear ($q = 0.138$). The 2022-to-2023 calendar-year changes point in the same direction: the language gap falls 24.6%, and the framework gap 36.9%.

5.5 Pretrends limit causal interpretation

The January 2023 coefficients are not evidence of a clean discontinuity from a stable pre-period. During January 2021–December 2022, Ruby share declines by 0.133 percentage points per month ($q < 0.001$), Rails by 0.104 points ($q = 0.003$), and Elixir by 0.024 points ($q = 0.065$). The language tar-

get/comparator gap declines 0.91% per month and the framework gap 1.18% per month; both have $q = 0.039$.

Figure 6 makes the problem visible. A model can estimate an additional January 2023 level difference while the outcome is already trending downward. In that setting, the level coefficient describes deviation from one fitted continuation; it does not identify the cause of the longer decline.

5.6 January 2023 is not a unique break date

The date-sensitivity analysis moves the assumed break from January 2021 through November 2023. For Ruby and Rails, 17.1% of candidate dates yield a fitted level change at least as negative as January 2023; the most negative dates are November 2022 and February 2023. For Elixir the fraction is 8.6%, with November 2022 most negative. The corresponding fractions are

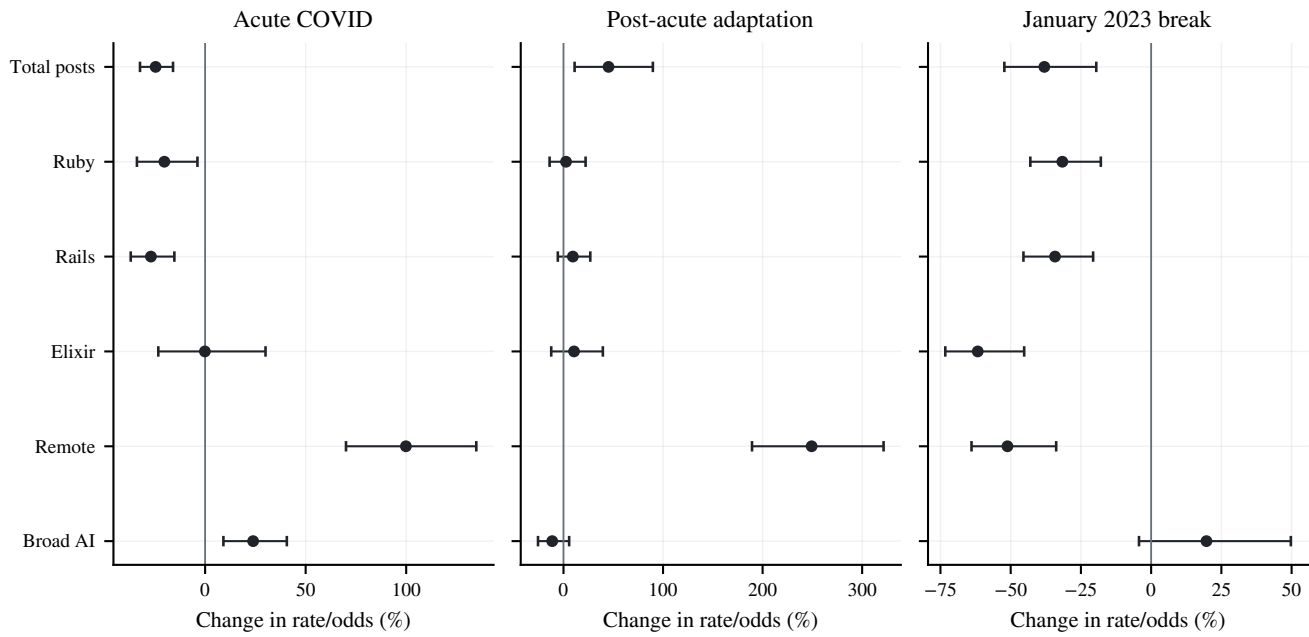


Figure 4. Primary model level effects. Points are incidence-rate changes for total posts and odds changes for term shares; bars are HAC(6) 95% confidence intervals. The three panels use different horizontal ranges. Effects are conditional temporal associations, not identified treatment effects.

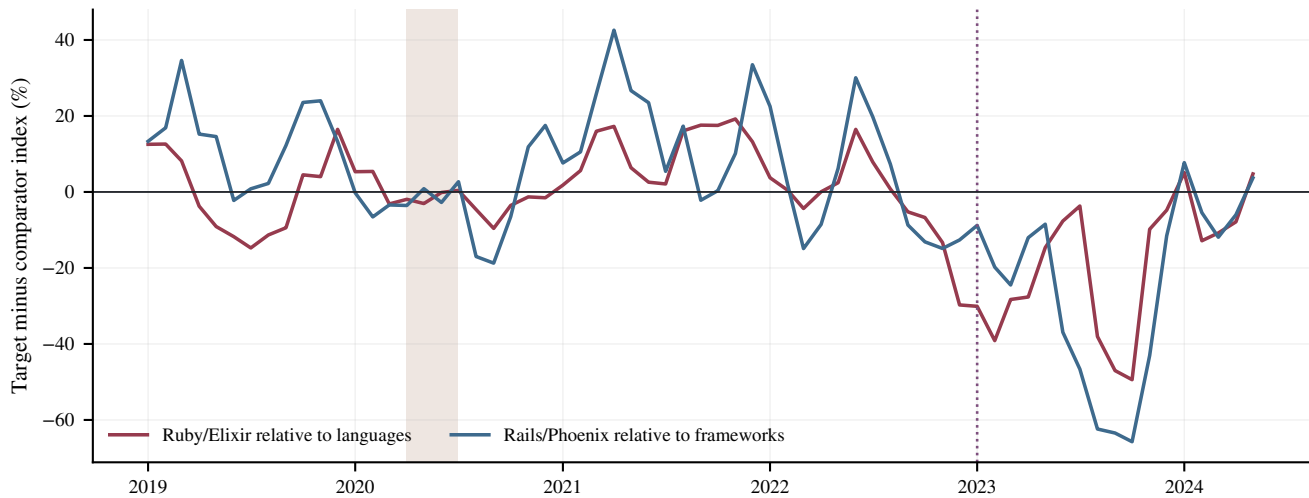


Figure 5. Target indices relative to comparator indices. Values are $100[\exp(g_t) - 1]$, where g_t is the log target/comparator index ratio. A value below zero indicates weaker target performance relative to the relevant comparator group.

Table 4. Comparator-adjusted January 2023 effects and pre-period trends.

Measure	Estimate	95% CI	<i>p</i>
Language gap: Jan. 2023	-35.9%	[-45.8, -24.2]	< 0.001
Framework gap: Jan. 2023	-38.0%	[-61.7, +0.3]	0.052
Ruby pretrend	-0.133 pp/mo	[-0.181, -0.086]	< 0.001
Rails pretrend	-0.104 pp/mo	[-0.167, -0.040]	0.001
Language-gap pretrend	-0.91%/mo	[-1.70, -0.11]	0.027
Framework-gap pretrend	-1.18%/mo	[-2.22, -0.13]	0.028

Notes: Gap effects compare normalized geometric-mean target and comparator indices. Pretrends are estimated over Jan. 2021–Dec. 2022 with HAC(3). Significant negative pretrends weaken a causal interpretation of a January 2023 discontinuity.

14.3% for the language gap and 31.4% for the framework gap, whose most negative candidate dates occur later in 2023.

These ranks do not invalidate the January 2023 event anchor. They show that the series does not present a uniquely timed

statistical break that can be cleanly assigned to ChatGPT’s launch. The relevant transition spans a wider market interval.

5.7 Within-era AI correlations are inconclusive

The consistent panel contains only 17 post-January-2023 months. Raw broad-AI correlations are -0.10 for Ruby, -0.36 for Rails, and $+0.01$ for Elixir. Raw LLM-intensity correlations are -0.14 , -0.36 , and $+0.10$. After separately removing linear time trends, all six correlations are small and none is significant; every adjusted q value is 0.963 (Table 5).

A null result in such a short series is not strong evidence of no relationship. It does, however, fail to support the specific mechanism that months with more AI or LLM terminology systematically have lower target-stack shares after accounting for linear time.

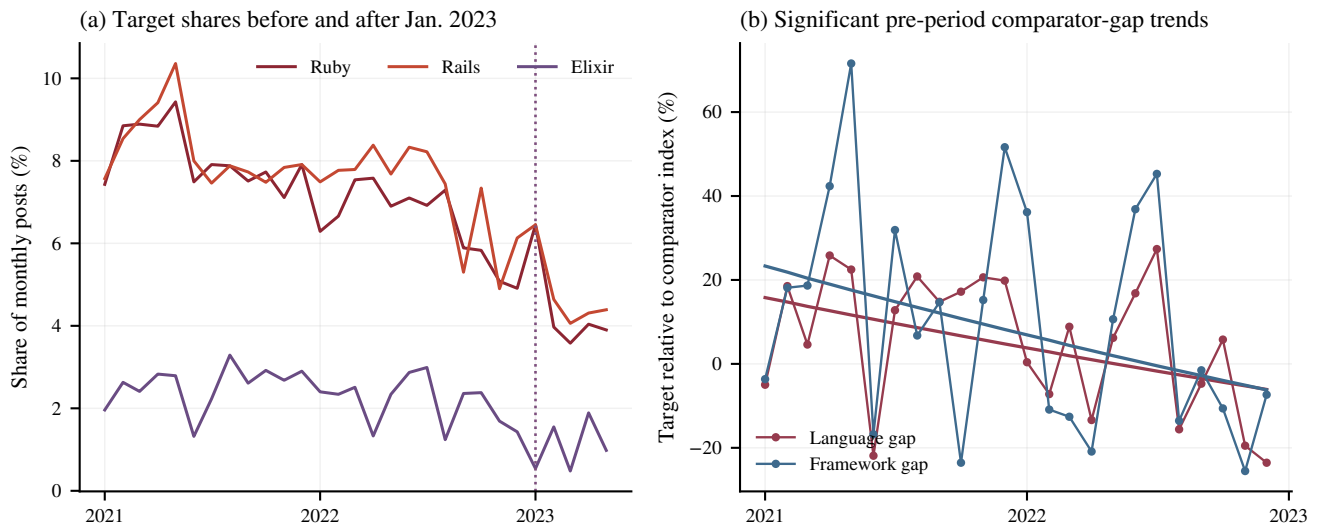


Figure 6. Pre-event target shares and comparator-gap trends. Panel (a) displays target shares from 2021 into the early post-launch period. Panel (b) shows the declining 2021–2022 comparator gaps and fitted pre-period trends. Significant pretrends are a central threat to causal attribution.

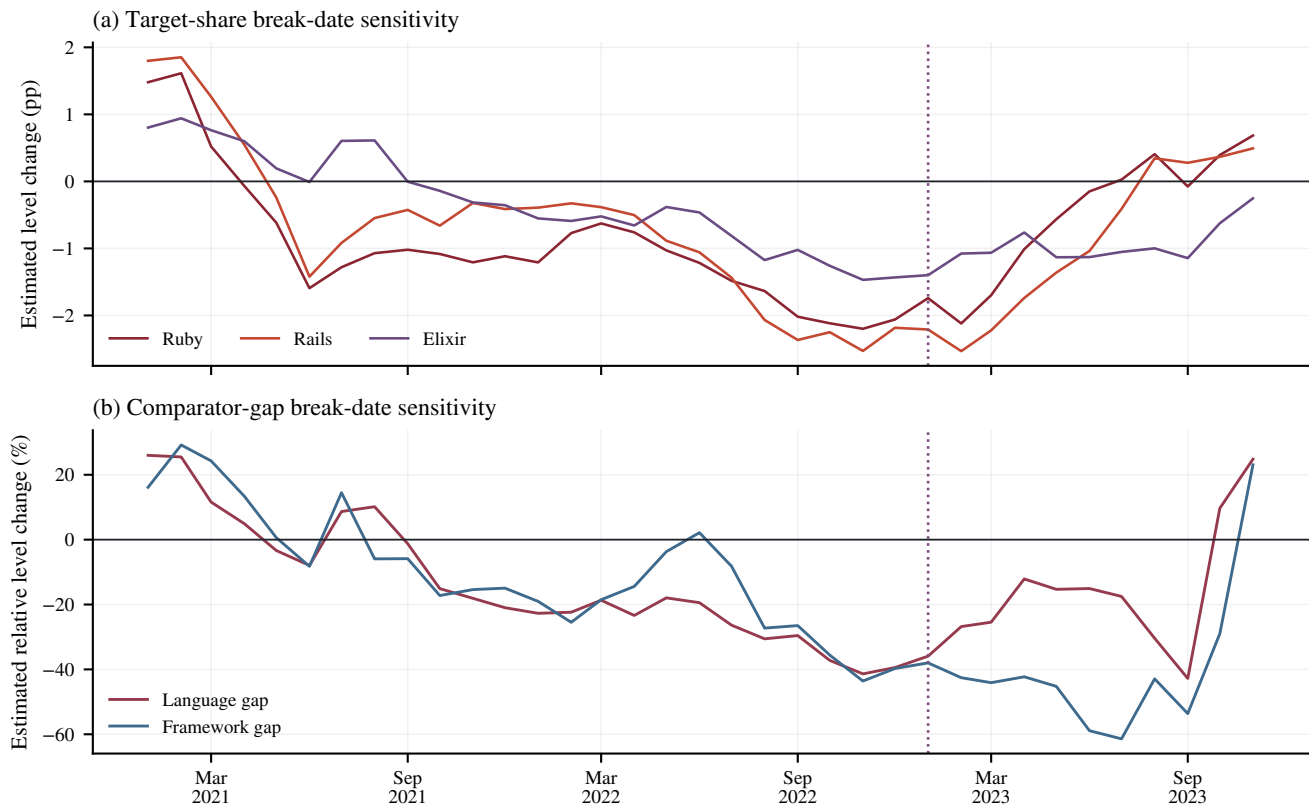


Figure 7. Break-date sensitivity. Each point is the fitted level coefficient when the January-2023 marker is replaced by the candidate month. The dotted line marks January 2023. Panel (a) is in percentage points; panel (b) converts log-gap coefficients to relative percent changes. This is a descriptive sensitivity plot, not a placebo randomization test.

5.8 Recent descriptive extension

The separate post-level extension shows stable total HN thread volume across three complete July–June periods but continuing decline in Ruby/Rails core share. Total posts are 3,822, 3,931, and 3,999, while the Ruby/Rails union falls from 5.44% to 4.66% to 4.10%. Elixir core counts are 45, 51, and 38. This pattern makes pure denominator contraction an incomplete

explanation for the latest Ruby/Rails decline, although it does not identify AI as the cause.

6 Robustness, Validity, and Ethics

6.1 What the models establish

The models establish conditional temporal associations in this corpus. The strongest one is the COVID-era remote-work shift. The January 2023 target-stack estimates are also robust in sign across direct and language-comparator specifications, but

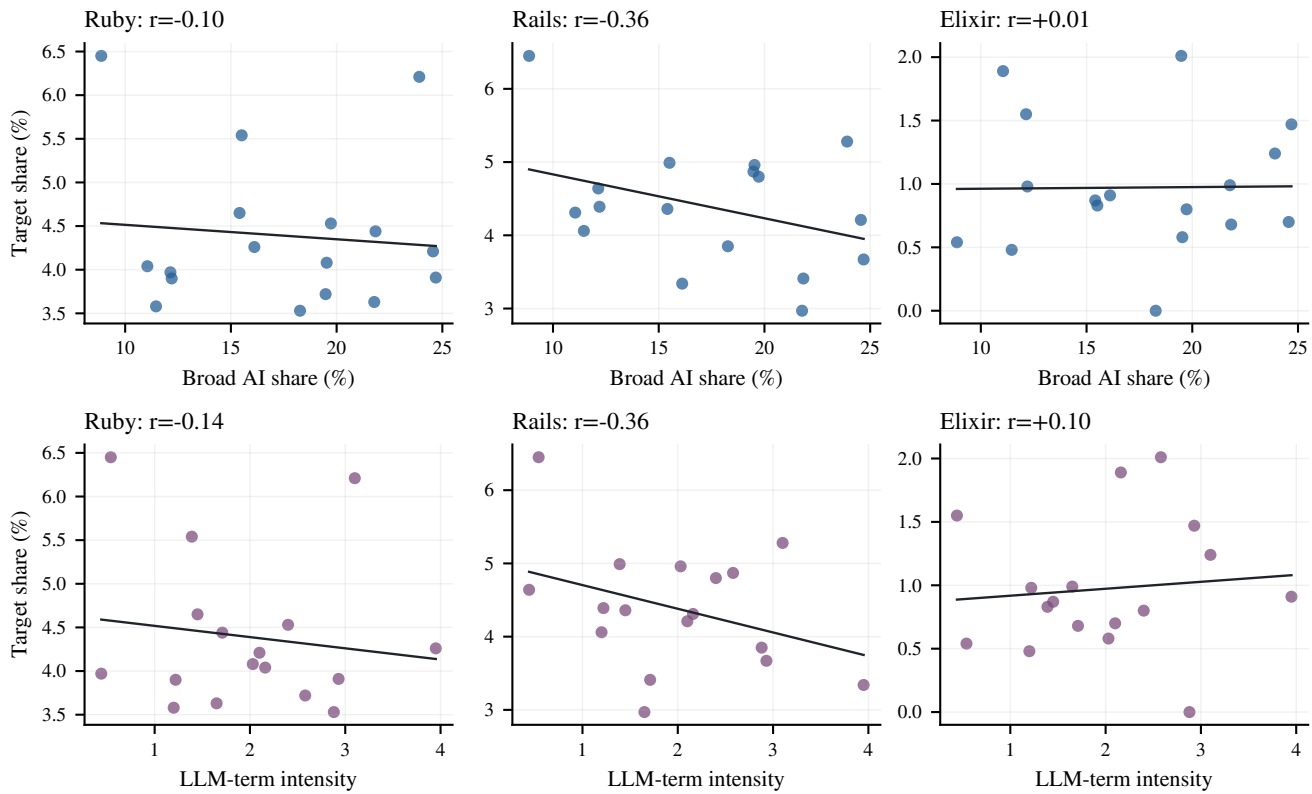


Figure 8. Raw monthly correlations during January 2023–May 2024. Top row: broad-AI share. Bottom row: additive LLM-term intensity. Lines are unadjusted least-squares fits. Detrended estimates are reported in Table 5.

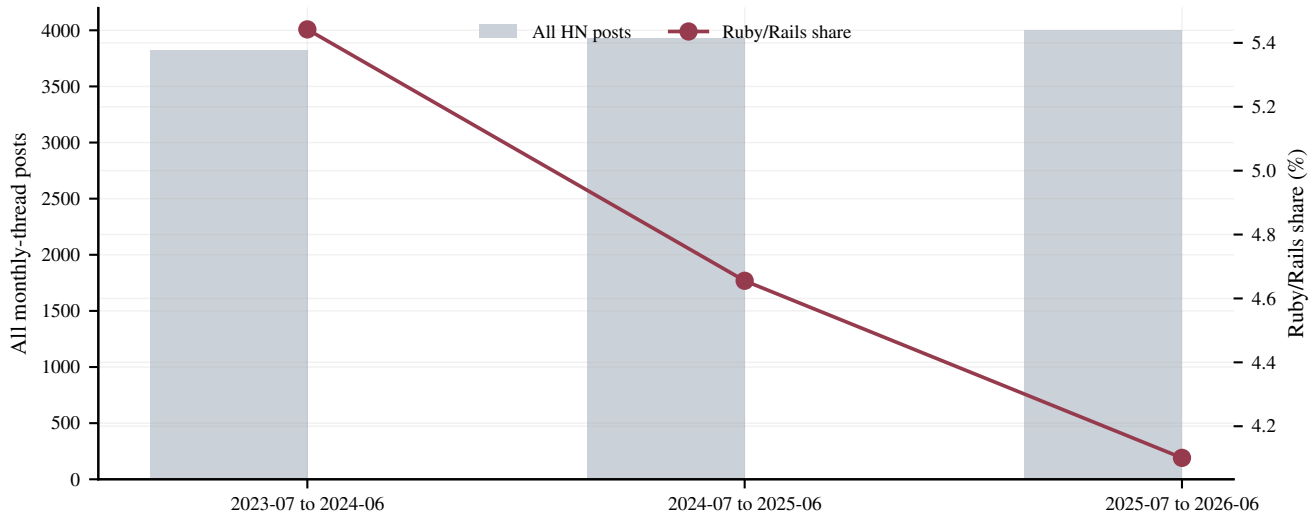


Figure 9. Recent complete July–June periods from the separately measured extension. Bars show total top-level posts across the corresponding monthly threads; the line shows the Ruby/Rails core share. The extension uses a different extraction pipeline and remains descriptive.

their causal interpretation is weak because of pretrends, nearly negative break dates, concurrent labor-market contraction, and absent firm-level treatment variation.

The term “post-ChatGPT period” should therefore be read chronologically. It does not mean that ChatGPT is the estimated treatment. A credible displacement design would require a unified post-level corpus, firm identifiers, adoption or exposure measures, comparison groups with defensible parallel trends, and a longer post-period.

6.2 Measurement and source transition

Aggregate agreement between HNTrends and the later curated extension is high (Figure 10), but correlation can coexist with systematic false positives and false negatives. Rails is especially vulnerable to uses such as payment or financial rails. Human precision and recall results are not reported because the named author has not yet completed the annotation protocol. A future API-based recall sample will also be limited to comments still retrievable at collection time, so deleted historical items must be reported as a survivorship limitation.

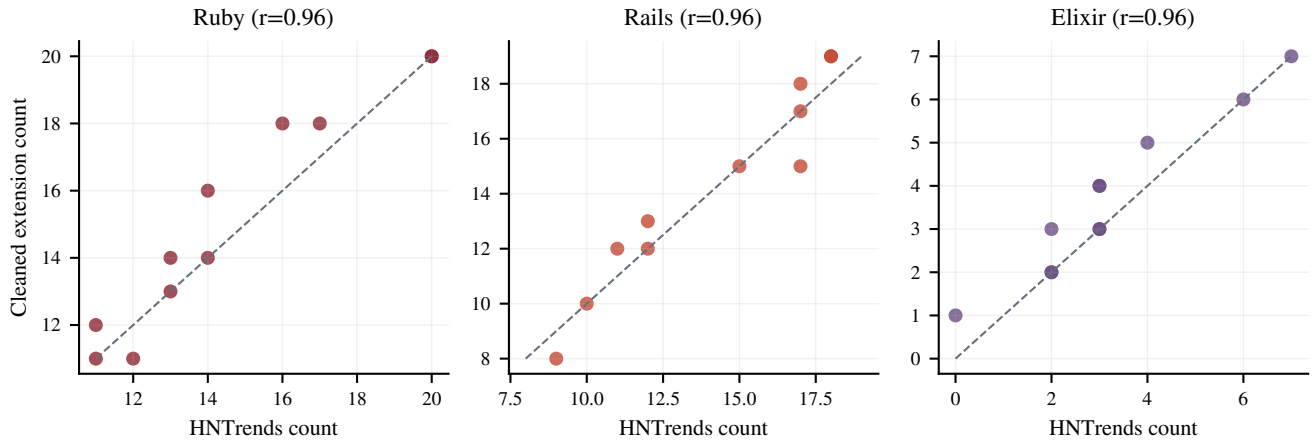


Figure 10. Monthly count agreement between HNTrends and the post-level extension during the 11-month overlap. The dashed line is equality. High aggregate correlation supports directional continuity but is not a post-level confusion matrix.

Table 5. Post-January 2023 correlation diagnostics.

Exposure	Outcome	r	FDR q
Broad AI share	Ruby	+0.24	0.963
Broad AI share	Rails	+0.15	0.963
Broad AI share	Elixir	-0.26	0.963
LLM-term intensity	Ruby	-0.05	0.963
LLM-term intensity	Rails	-0.17	0.963
LLM-term intensity	Elixir	+0.05	0.963

Notes: January 2023–May 2024 ($n = 17$ months). Both exposure and outcome are residualized on separate linear time trends. The LLM measure is additive term intensity, not a unique-post share.

Table 6. Recent complete July–June periods.

Period	All posts	Ruby/Rails	Share	Elixir
2023-07 to 2024-06	3,822	208	5.44%	45
2024-07 to 2025-06	3,931	183	4.66%	51
2025-07 to 2026-06	3,999	164	4.10%	38

Notes: The post-level extraction differs from the formal HNTrends panel. July 2026 is partial and excluded. Ruby/Rails and Elixir counts can overlap.

6.3 Selection and external validity

HN is a niche community with a particular employer mix, technical culture, and geographic distribution. Changes can reflect which firms choose to post, not only changes in vacancies within firms. A single comment may contain multiple roles. Repeated postings may represent persistent openings, refreshed advertisements, or separate hiring rounds. The results should not be generalized to the entire software labor market without external comparison.

6.4 Privacy and data minimization

The public post-level release excludes raw text, usernames, contact emails, external application URLs, and target-context excerpts. Canonical HN item IDs and URLs remain for source verification. Automated scans of the sanitized object columns found zero residual email-pattern and zero non-HN URL-pattern hits. The package does not claim ownership of HN user posts and separates licenses for original code, author-created material, derived aggregate outputs, and third-party sources.

6.5 Generative-AI assistance and authorship

OpenAI ChatGPT was used extensively for analysis code, statistical workflow, drafting, figures, documentation, typesetting, and QA. It is not an author. ACM policy, for example, prohibits generative-AI tools from authorship and requires disclosure of

substantive generated content [19]. Marian Posăceanu is the sole named author and must verify the data, code, citations, numerical results, and prose before submission. The included sign-off checklist, blinded validation packet, separate scoring key, and recall-sampling helper make the remaining human quality gates explicit.

7 Discussion

Three interpretations are consistent with the evidence.

First, COVID-19 altered the mode of software recruitment more clearly than the relative visibility of the target stacks. Remote terminology more than doubled in odds during the acute period and increased further during adaptation; comparator-adjusted target penalties around COVID are statistically uncertain. This aligns with research describing a rapid and persistent transition toward remote-capable work, including heterogeneous experiences within software teams [8–10].

Second, Ruby, Rails, and Elixir occupy a smaller portion of the HN hiring discourse after 2022. The language target index declines relative to a broad language comparator set, and the recent extension shows a continued Ruby/Rails share decrease even as total thread volume stabilizes. This is evidence of relative visibility loss in this channel.

Third, the data do not identify LLM displacement as the mechanism. Broad AI terminology accelerates, but broad AI is not LLM adoption. LLM-specific intensity is not significantly correlated with target shares in the short post-launch window. More importantly, Ruby, Rails, and comparator gaps were already declining before 2023, and the estimated break is not unique to January 2023. A broader market reallocation, changing startup composition, technology substitution unrelated to LLMs, financing conditions, or longer lifecycle trends remain viable explanations.

The distinction matters because experimental productivity evidence and vacancy-demand evidence answer different questions. A tool can increase productivity for particular tasks while employment rises, falls, or changes composition depending on output demand, organizational redesign, and complementary investment [11–13]. This HN corpus cannot resolve those margins.

Table 7. Release-quality gates and current status.

Gate	Status	Evidence or required action
Authorship	Corrected	Sole named author is Marian Posăceanu; no model or vendor is listed as an author.
Privacy	Passed for public CSVs	Raw text, usernames, contact-email fields, target excerpts, and application URLs removed. Automated scan found 0 residual email-pattern hits.
Precision validation	Pending author annotation	Complete the 225-row blinded packet, then score it with the separate weighted prediction key.
Recall validation	Not yet estimable	Draw and annotate a probability sample from all monthly top-level posts; the scorer reports recall only for that design.
Causal identification	Reframed	Manuscript reports temporal associations; significant 2021–2022 pretrends and break-date sensitivity are presented explicitly.
Reproducibility	Implemented	Selected aggregate panel, source checksums/provenance, sanitized derived data, code, generated outputs, build script, and SHA-256 manifest are included.

8 Conclusion

The monthly HN hiring archive records two different kinds of change. Around COVID-19, total posting volume contracts acutely and remote-work terminology rises sharply and persistently. Around the post-ChatGPT period, target-stack visibility is lower and broad-AI language accelerates. The latter concurrence is not sufficient to attribute Ruby, Rails, or Elixir decline to LLMs.

The revised analysis improves the evidentiary basis through appropriate count/share models, comparator ecosystems, HAC uncertainty, FDR correction, pretrend tests, break-date sensitivity, and a separate recent extension. Those changes make the negative conclusion more important, not less: the available data support temporal association and relative decline, but not direct causal displacement.

Before peer-reviewed submission, the author should complete post-level human precision and recall validation, independently rerun the analysis, verify every citation and numerical statement, and adapt the artifact to the target venue's current policy. Until then, the appropriate status is an exploratory, reproducible working paper with explicit unresolved quality gates.

Data and Code Availability

The accompanying public package contains the selected aggregate analytical panel, source checksums and provenance, sanitized derived post-level datasets, generated tables and figures, analysis and validation scripts, LaTeX source, privacy report, documentation, and SHA-256 manifest. The upstream HNTrends snapshot, raw HN comment text, direct contact fields, and API caches are excluded. Human annotation outputs are not included because annotation remains pending.

Acknowledgments and AI Disclosure

OpenAI ChatGPT (GPT-5.6 Pro, accessed July 2026) assisted with data-processing and statistical code, manuscript drafting and revision, figure production, documentation, typesetting, and quality checks. The named author must review and verify all outputs before submission and accepts responsibility only after completing the bundled sign-off process. No model or vendor is listed as an author.

References

- [1] D. Deming and L. B. Kahn, "Skill requirements across firms and labor markets: Evidence from job postings for professionals," *Journal of Labor Economics*, vol. 36, no. S1, pp. S337–S369, 2018. doi: 10.1086/694106.
- [2] B. Hershbein and L. B. Kahn, "Do recessions accelerate routine-biased technological change? Evidence from vacancy postings," *American Economic Review*, vol. 108, no. 7, pp. 1737–1772, 2018. doi: 10.1257/aer.20161570.
- [3] D. Acemoglu, D. Autor, J. Hazell, and P. Restrepo, "Artificial intelligence and jobs: Evidence from online vacancies," *Journal of Labor Economics*, vol. 40, no. S1, pp. S293–S340, 2022. doi: 10.1086/718327.
- [4] HNTrends, "Hacker News Hiring Trends," aggregate monthly trend archive, 2026. hntrends.com.
- [5] World Health Organization, "WHO Director-General's opening remarks at the media briefing on COVID-19—11 March 2020," 2020. WHO briefing.
- [6] OpenAI, "Introducing ChatGPT," 30 Nov. 2022. OpenAI release page.
- [7] J. I. Dingel and B. Neiman, "How many jobs can be done at home?" *Journal of Public Economics*, vol. 189, art. 104235, 2020. doi: 10.1016/j.jpubeco.2020.104235.
- [8] J. M. Barrero, N. Bloom, and S. J. Davis, "The evolution of work from home," *Journal of Economic Perspectives*, vol. 37, no. 4, pp. 23–50, 2023. doi: 10.1257/jep.37.4.23.
- [9] D. Ford, M.-A. Storey, T. Zimmermann, C. Bird, S. Jaffe, C. Maddila, J. L. Butler, B. Houck, and N. Nagappan, "A tale of two cities: Software developers working from home during the COVID-19 pandemic," *ACM Transactions on Software Engineering and Methodology*, vol. 31, no. 2, art. 27, pp. 1–37, 2022. doi: 10.1145/3487567.
- [10] D. Russo, P. P. H. Hanel, S. Altnickel, and N. van Berkel, "The daily life of software engineers during the COVID-19 pandemic," in *Proc. IEEE/ACM ICSE-SEIP*, pp. 364–373, 2021. doi: 10.1109/ICSE-SEIP52600.2021.00048.
- [11] S. Noy and W. Zhang, "Experimental evidence on the productivity effects of generative artificial intelligence," *Science*, vol. 381, no. 6654, pp. 187–192, 2023. doi: 10.1126/science.adh2586.
- [12] E. Brynjolfsson, D. Li, and L. Raymond, "Generative AI at work," *Quarterly Journal of Economics*, vol. 140, no. 2, pp. 889–942, 2025. doi: 10.1093/qje/qjae044.
- [13] S. Peng, E. Kalliamvakou, P. Cihon, and M. Demirer, "The impact of AI on developer productivity: Evidence from GitHub Copilot," arXiv:2302.06590, 2023. doi: 10.48550/arXiv.2302.06590.
- [14] J. Lopez Bernal, S. Cummins, and A. Gasparrini, "Interrupted time series regression for the evaluation of public health interventions: A tutorial," *International Journal of Epidemiology*, vol. 46, no. 1, pp. 348–355, 2017. doi: 10.1093/ije/dyw098. See also the published corrigendum.
- [15] W. K. Newey and K. D. West, "A simple, positive semi-definite, heteroskedasticity and autocorrelation consistent covariance matrix," *Econometrica*, vol. 55, no. 3, pp. 703–708, 1987. doi: 10.2307/1913610.
- [16] Y. Benjamini and Y. Hochberg, "Controlling the false discovery rate: A practical and powerful approach to multiple testing," *Journal of the Royal Statistical Society: Series B*, vol. 57, no. 1, pp. 289–300, 1995. doi: 10.1111/j.2517-6161.1995.tb02031.x.
- [17] *hacker-job*, "Searchable jobs and hiring trends from Hacker News Ask HN: Who is Hiring? threads," source repository, 2026. GitHub repository.
- [18] Hacker News, "API documentation," official Firebase API reference, accessed 19 July 2026. github.com/HackerNews/API.
- [19] Association for Computing Machinery, "Policy on authorship" and generative-AI frequently asked questions, 2026. Authorship policy; FAQ.