

Pivoting or Leaping?

How the Structure of Exploration Shapes the Reward of Atypical Combination

Saeyoung Choi, Junghye Lee

Seoul National University

2026 International Conference on the Science of Science and Innovation

This work was supported by the NRF Grant funded by the Korea Government (MSIT) under Grant No. RS-2025-00516776. This work was also supported by the Ministry of Education and the NRF of Korea (NRF-2020S1A3A2A02093277).

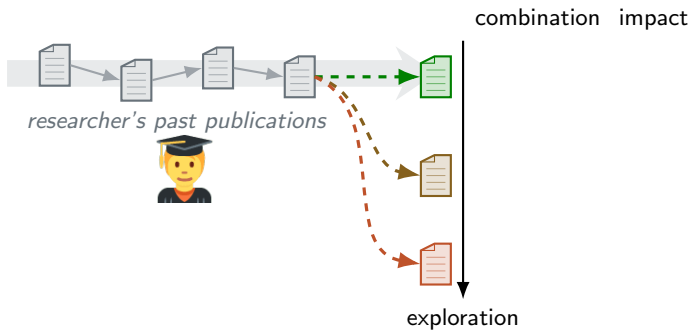
The exploration trade-off in scientific research

What happens when scientists explore beyond their prior expertise?

What happens when scientists explore beyond their prior expertise?

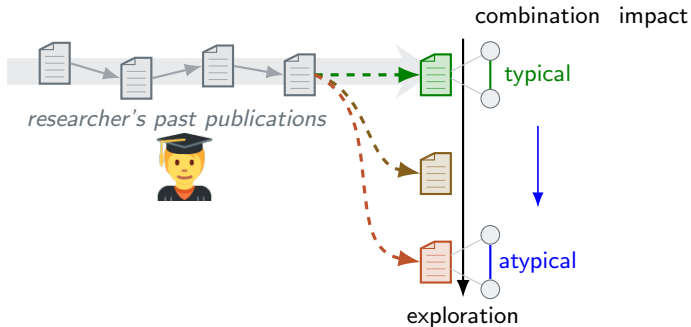
- *Arts & Fleming (2018)*: When inventors explore, *novelty increases* but *value decreases*.
 - Negative effect on value is muted when collaborating or using scientific knowledge.
- *Hill et al. (2025)*: The further researchers explore, *the lower the impact*.
 - High-pivot work was associated with *a higher propensity for atypical combinations*.

Two opposite influences of exploration



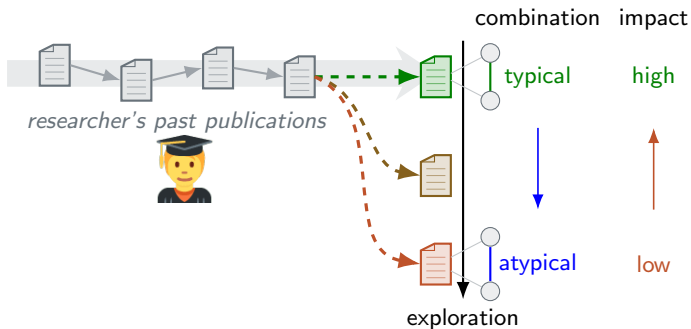
Two opposite influences of exploration

- Scientists must leave familiar ground to find atypical combinations,



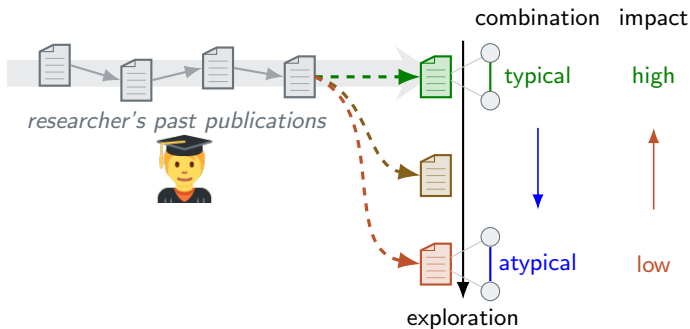
Two opposite influences of exploration

- Scientists must leave familiar ground to find atypical combinations, yet the further they explore, the lower the impact.



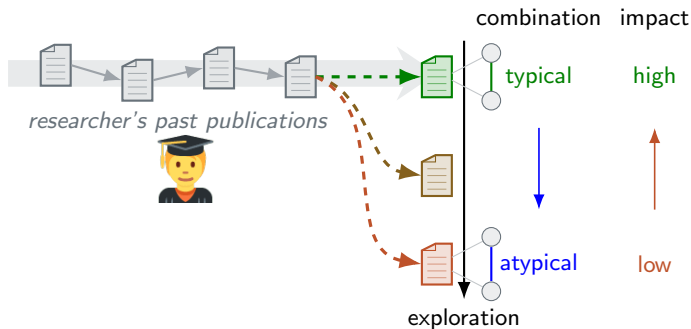
Two opposite influences of exploration

- Scientists must leave familiar ground to **find atypical combinations**, yet the further they explore, **the lower the impact**.
- So it looks like a choice: **stay safe and conventional**, or **go novel and risky**.



Two opposite influences of exploration

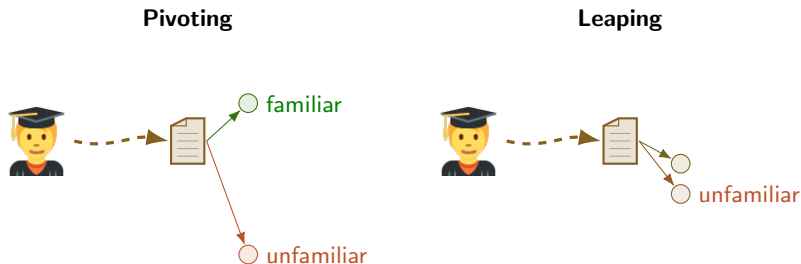
- Scientists must leave familiar ground to **find atypical combinations**, yet the further they explore, **the lower the impact**.
- So it looks like a choice: **stay safe and conventional**, or **go novel and risky**.



But this assumes exploration is **a single dimension**.

Pivoting or leaping?

- An **atypical combination** joins two components rarely cited together (globally unusual).
- But how **familiar** is each component to *the researcher*? Does they **pivot** from a familiar **anchor** into the unfamiliar, or **leap** without any anchor?



Pivoting or leaping — does it change what becomes atypical, and what becomes impactful?

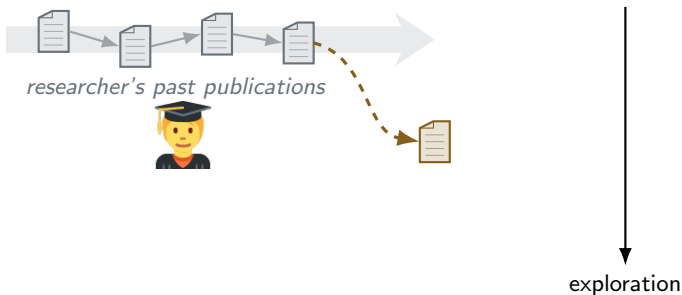
Pivoting or leaping — does it change what becomes atypical, and what becomes impactful?

- **RQ1** (*structure*) — Is atypical combination produced by relying on familiar knowledge, or by leaping into unfamiliar knowledge?
- **RQ2** (*impact*) — Is an atypical combination rewarded more when it is anchored (a pivot) than when it is unanchored (a leap)?

A granular view of exploration

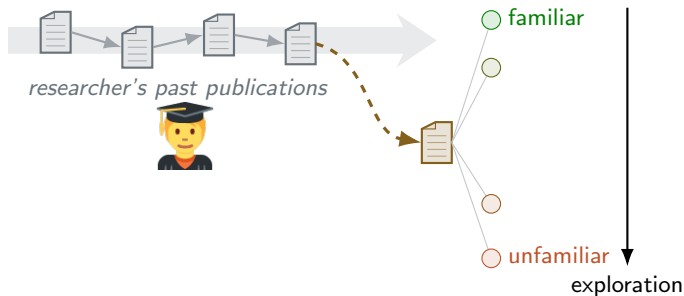
A granular view of exploration

- Instead of measuring one-dimensional exploration distance, we measure:

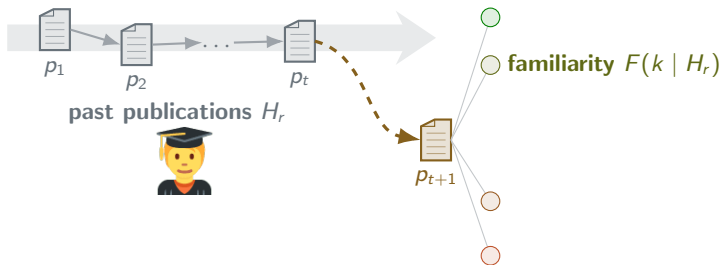


A granular view of exploration

- Instead of measuring one-dimensional exploration distance, we measure:
- researchers' **familiarity** with each **knowledge component**.
 - Knowledge component: knowledge unit used to construct new combinations, e.g., what kind of journals are cited in a paper.



Measuring component familiarity



- Component familiarity reflects **how likely a researcher is to use** a component k given the researcher's prior publication history H_r .
- It should
 - be **conditioned** on the researcher's prior trajectory.
 - rank even **out-of-trajectory** components, not just those that were observed in the trajectory.

Measuring component familiarity

- Component familiarity is *not* a new idea — but prior measures do not satisfy the two requirements.

	Component-level	Trajectory-conditioned	Out-of-trajectory
Pivot size (Hill et al. 2025)	× (paper-level)	○ (distance)	×
Component familiarity (Fleming 2001)	○	× (population popularity)	○
Search depth & scope (Ahuja & Katila 2001)	○	○ (repetition & exploration)	×

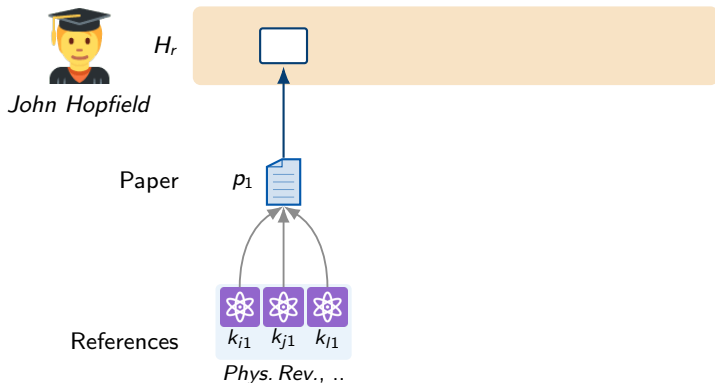
To overcome these limitations,
we have to estimate the conditional probability, $\Pr(k \mid H_r)$.

Estimating $\Pr(k \mid H_r)$ with a *sequential recommendation* model

- A sequential model predicts *the journals a researcher is likely to cite in the next paper*, conditional on their prior publication history — thus estimating $\Pr(k \mid H_r)$.

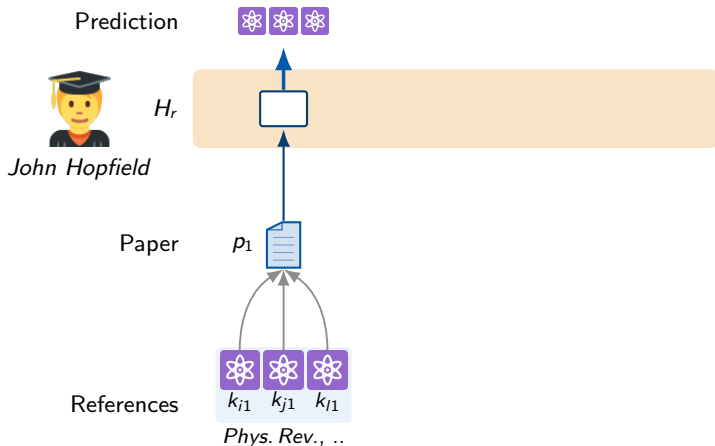
Estimating $\Pr(k \mid H_r)$ with a *sequential recommendation* model

- A sequential model predicts *the journals a researcher is likely to cite in the next paper*, conditional on their prior publication history — thus estimating $\Pr(k \mid H_r)$.



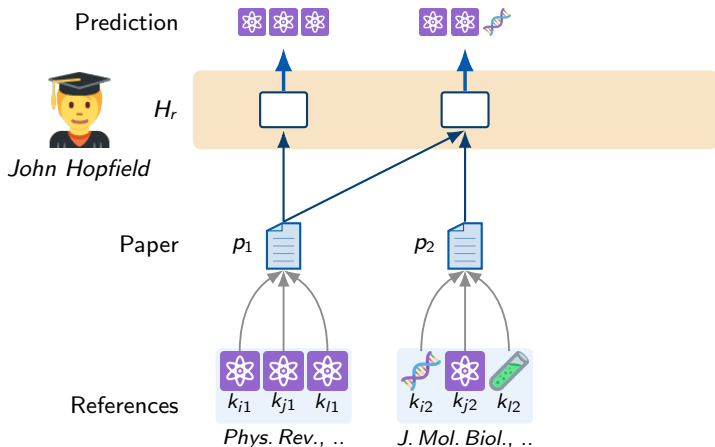
Estimating $\Pr(k \mid H_r)$ with a *sequential recommendation* model

- A sequential model predicts *the journals a researcher is likely to cite in the next paper*, conditional on their prior publication history — thus estimating $\Pr(k \mid H_r)$.



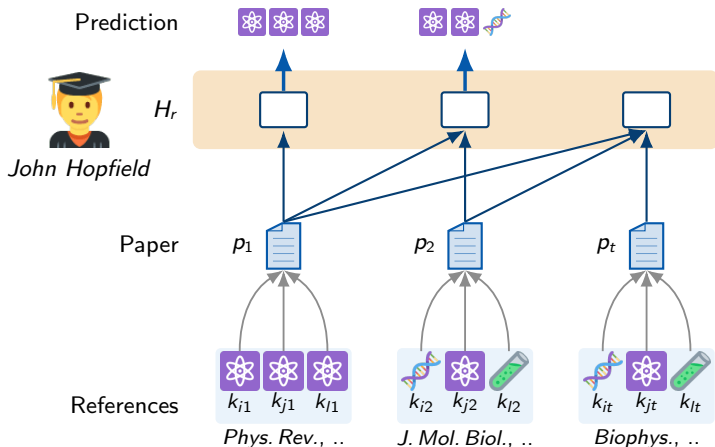
Estimating $\Pr(k \mid H_r)$ with a *sequential recommendation* model

- A sequential model predicts *the journals a researcher is likely to cite in the next paper*, conditional on their prior publication history — thus estimating $\Pr(k \mid H_r)$.



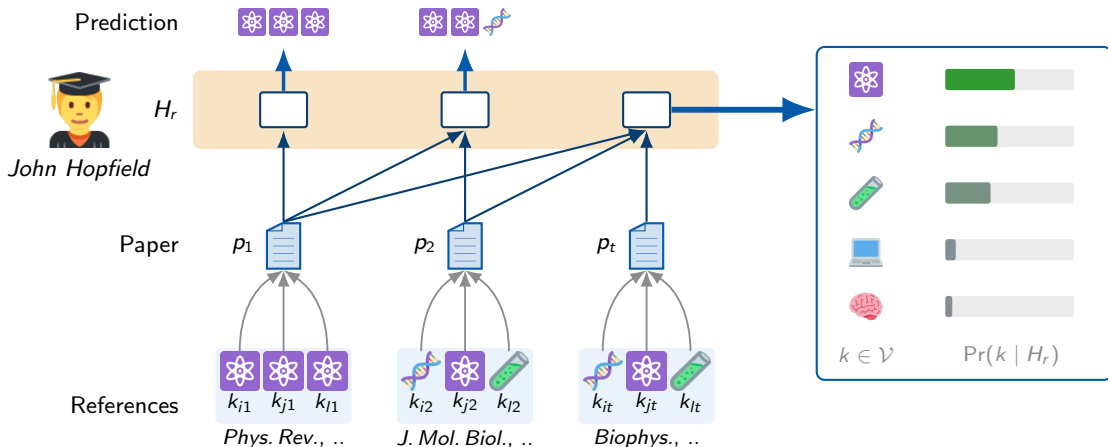
Estimating $\Pr(k \mid H_r)$ with a *sequential recommendation* model

- A sequential model predicts *the journals a researcher is likely to cite in the next paper*, conditional on their prior publication history — thus estimating $\Pr(k \mid H_r)$.



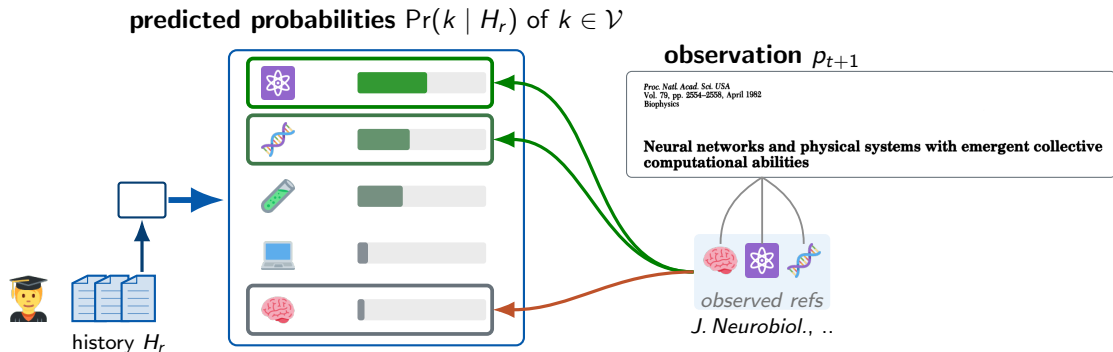
Estimating $\Pr(k \mid H_r)$ with a *sequential recommendation* model

- A sequential model predicts *the journals a researcher is likely to cite in the next paper*, conditional on their prior publication history — thus estimating $\Pr(k \mid H_r)$.



Predicted probabilities of observations

- We then compare what the model predicts with what actually happens.



- If an observed reference has a **high predicted probability**, it is **familiar** to the researcher; if it has a **low predicted probability**, it is **unfamiliar**.

Metric definitions

- **Component familiarity** (proposed)

$$\tilde{F}(k | H_r) = \frac{F(k | H_r) - \bar{F}}{\bar{F}_{\text{repeat}} - \bar{F}}, \quad F(k | H_r) = \log \Pr(k | H_r) - \log \Pr(k)$$

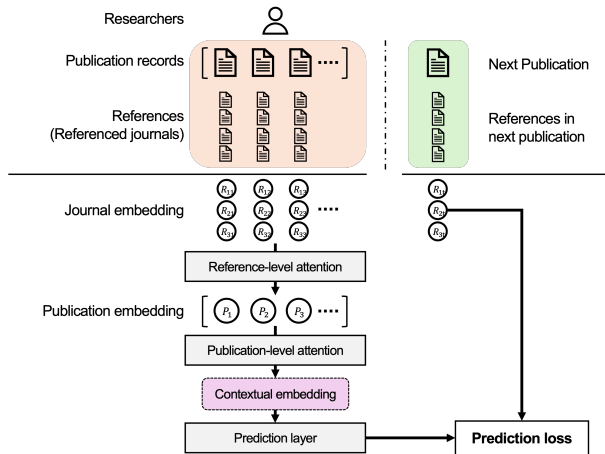
- F represents the log-odds of observing k given the researcher's prior trajectory, relative to the global popularity of k .
- $\tilde{F}$ is normalized with respect to the researcher's prior trajectory: $\bar{F}$ is the average F over the vocabulary, and $\bar{F}_{\text{repeat}}$ is the average F over the researcher's prior trajectory.
- High value means the component is relatively familiar to the researcher; low value means the component is relatively unfamiliar to the researcher.

- **Combination atypicality** (Uzzi et al., 2013, Lin et al., 2022)

$$A(k_i, k_j) = -\text{npmi}(k_i, k_j) = -\frac{\log(\Pr(k_i) \Pr(k_j)) - \log \Pr(k_i, k_j)}{-\log \Pr(k_i, k_j)}$$

- representing how unusual a combination of components is: low value means the pair is commonly observed together; high value means the pair is rarely observed together.

Backup: sequential model details



Hierarchical self-attentive model

- Journal (reference) $e_k = \text{Emb}(k)$
- Paper $p_t = \text{Attn}(\{e_k\}_{k \in p_t})$
- Researcher $h_t = \text{Causal_Attn}(x_{1:t})$
- $\Pr(k | H_r) = \text{softmax}_k(\text{score}(h_t, e_k))$
- Weighted cross-entropy loss

$$\mathcal{L} = \sum_{k \in P_{t+1}} \frac{c_k}{\sum_c} [-\log \Pr(k | H_r)]$$

Specification

- Journal embedding dim 2048; 1 layer; year + career-year embeddings.
- ≤ 64 recent papers for each researcher;
 ≤ 128 references for each paper.

Results

Model performance

- Heuristic baselines
 - Popularity: predicting the most globally popular journals
 - Repetition: predicting the most frequently cited journals in the researcher's prior trajectory

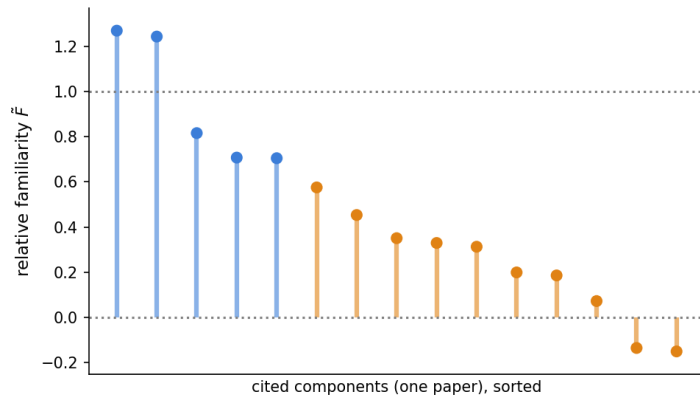
Method	Recall@50			nDCG@50
	ALL	Seen	Unseen	
Popularity	0.122	0.169	0.073	0.093
Repetition (5 papers)	0.422	0.835	0.000	0.479
Repetition (all)	0.427	0.844	0.000	0.476
Sequential model	0.468	0.752	0.178	0.493

Next-paper reference journal prediction for 2020 (model trained on 2010–2019).

Trends are similar for other years (1980–2024).

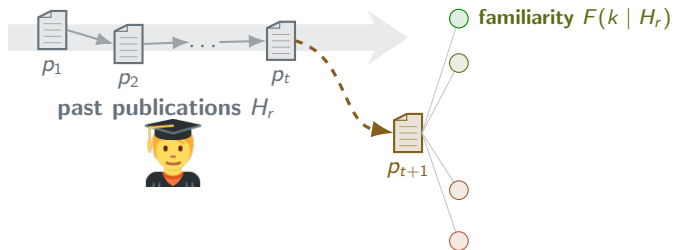
- Repeat heuristics perform well for **already seen** journals but **can't predict unseen** ones.
- The proposed sequential model outperforms baselines for overall scores, as it can predict researchers' exploration into **unseen** journals.

Distribution of component familiarity

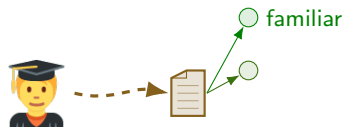


- A single paper mixes **anchors** (high $\tilde{F}$) with **unfamiliar** ones (low $\tilde{F}$).

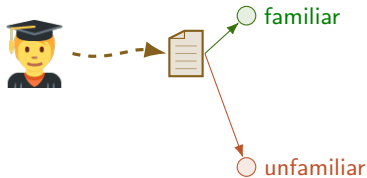
Pair-level structure: component familiarity & combination atypicality



Routine



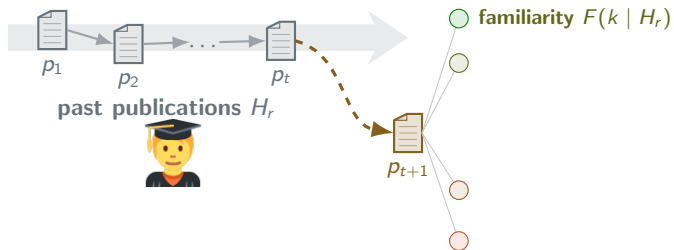
Pivoting



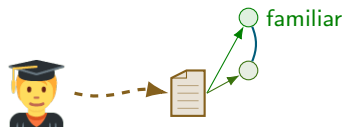
Leaping



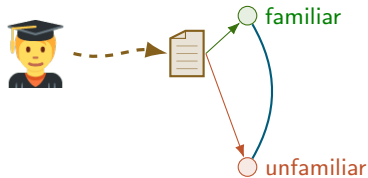
Pair-level structure: component familiarity & combination atypicality



Routine



Pivoting



Leaping

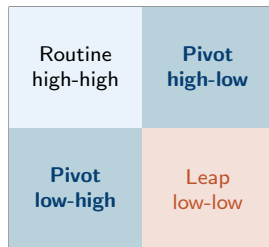
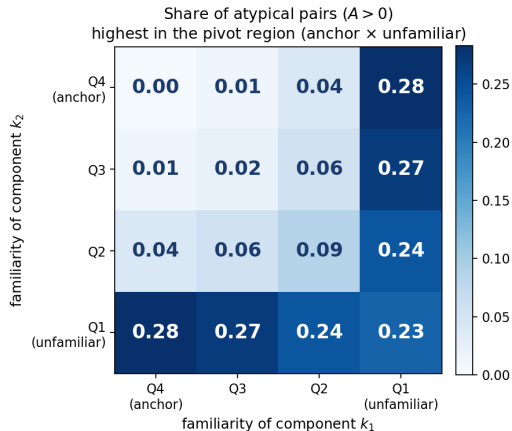


Which pair types produce atypical combinations?

- **Is atypical combination produced by relying on familiar knowledge, or by leaping into unfamiliar knowledge?**

Which pair types produce atypical combinations?

- Is atypical combination produced by relying on familiar knowledge, or by leaping into unfamiliar knowledge?



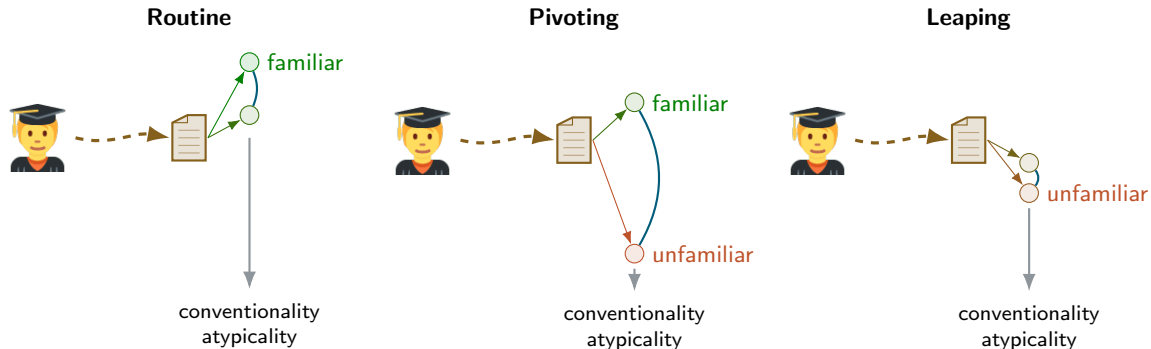
- Atypicality is highest in the asymmetric region (familiar $\times$ unfamiliar quartile).

From pair-level structure to paper-level outcomes

- Is an atypical combination rewarded more when it is anchored (a **pivot**) than when it is unanchored (a **leap**)?

From pair-level structure to paper-level outcomes

- Is an **atypical combination** rewarded more when it is anchored (a **pivot**) than when it is unanchored (a **leap**)?
- We model a paper's impact from the **conventionality** and **atypicality** of its reference pairs, estimated *separately within each pair type* (routine / **pivot** / **leap**):



From pair-level structure to paper-level outcomes

- Is an **atypical combination** rewarded more when it is anchored (a **pivot**) than when it is unanchored (a **leap**)?
- We model a paper's impact from the **conventionality** and **atypicality** of its reference pairs, estimated *separately within each pair type* (routine / **pivot** / **leap**):

$$\text{logit Pr}(\text{top-5}\%_p) = \beta_0 + \sum_{c \in \{\text{rt}, \text{pv}, \text{lp}\}} (\beta_c \text{Conv}_p^c + \gamma_c \text{Atyp}_p^c) + \mathbf{X}_p' \delta$$

Predictor	Definition
Conventionality Conv_p within routine / pivot / leap	conventionality of reference pairs ($-\text{med } A$). conventionality, computed <i>within</i> each pair type.
Share of atypical pairs Atyp_p within routine / pivot / leap	share of atypical pairs within reference pairs ($A > 0$). the same share, computed <i>within</i> each pair type.
Controls ($\mathbf{X}_p$)	share of familiar components, pivot size (cosine distance between p and H_r), team size, number of references, career age, prior track record, and year fixed effects.

Logit on top-5% impact (within field-year); 1980–2020, with ≥ 5 prior papers.

The value of atypicality depends on anchoring

	Uzzi et al. (2013)		decomposition
	Model (1)	Model (2)	Model (3)
Conventionality	+2.96***	+2.20***	
routine			+1.06***
pivot			+1.44***
leap			−0.31***
Tail atypicality (90th percentile)	+1.70***		
Share of atypical pairs		+1.40***	
routine			+0.96***
pivot			+0.49***
leap			+0.26***
observations	4,405,798	4,405,798	4,405,798
pseudo- R^2	0.129	0.128	0.130

Logit on top-5% impact; 1980–2020; controls as listed; *** $p < 0.01$.

- **Conventionality** pays *only when anchored*; an unanchored conventionality does not help.
- Atypicality increases the impact, but the **anchoring** matters — routine > pivot > leap ($p < 0.001$ for all three pairwise differences).

Distinguishing pivot from leap: the anchoring premium

- Our **component-level** familiarity distinguishes a **pivot** from a **leap**.
- This structure **interacts** with combination atypicality.

Implication

What makes atypical work pay is whether it stays **anchored** — the **anchoring premium**. Evaluating science by exploration or combination alone misses this nuanced structure.

The anchoring premium shrinks as teams grow

	team size		
	1 (solo)	2	3+
Conventionality			
routine	+1.06***	+1.14***	+0.98***
pivot	+1.87***	+1.48***	+1.02***
leap	-0.63***	-0.28***	-0.31***
Share of atypical pairs			
routine	+1.34***	+0.96***	+0.78***
pivot	+0.87***	+0.63***	+0.31***
leap	-0.09*	+0.24***	+0.26***
observations	577,879	861,284	2,966,635
pseudo- R^2	0.153	0.146	0.119

Logit on top-5% impact; *** $p < 0.01$, * $p < 0.1$.

- The anchoring premium is steep for **solo** (**pivot** $\gg$ **leap**); as teams grow it flattens.
- In **3+ teams**, they are **not significantly different** ($p = 0.24$).

Teams turn leaps into pivots (ongoing work)



- A component unfamiliar to one author can be an **anchor** for a co-author.
- Teams enable the anchored exploration that cannot be achieved by individual researchers.

- **What we did:**

- trained a **sequential deep learning model** to predict which knowledge components a researcher will cite next, based on their prior trajectory.
- estimated each researcher's **component familiarity** from a sequential model.
- defined three pair types based on the components' familiarity — **routine** / **pivot** / **leap**.

- **Findings:**

- **Structure:** combination atypicality is highest in the asymmetric region, with a familiar anchor reaching a distant component (**pivot**).
- **Impact:** the atypicality premium is conditional on anchoring — routine > **pivot** > **leap**.
- **Teams** (ongoing): the **pivot-leap** premium fades in larger teams, as collaborators can supply anchors for each other.

Thank you for your attention!
Questions and feedback are very welcome.

saeyoung.choi@snu.ac.kr

Backup: prior work

- Title — the word “anchored exploration”
 - *Park, M., Maity, S. K., Wuchty, S., & Wang, D. (2026). Interdisciplinary papers supported by disciplinary grants garner deep and broad scientific impact. PNAS nexus, 5(3), pgag057.*
 - ‘To unleash the full potential of interdisciplinary research, funders and policymakers might consider mechanisms that encourage **“anchored exploration”** rather than enforcing artificial recombination of distant disciplines.’
- Exploration in scientific research
 - *Arts, S., & Fleming, L. (2018). Paradise of novelty—or loss of human capital? Exploring new fields and inventive output. Organization Science, 29(6), 1074-1092.*
 - **Explore**: no overlapping patent classes between focal patent and prior patents.
 - **Novelty**: (1) the number of *new words* appearing in a patent (title, abstract, claims) that have not appeared in prior patents and (2) *backward citations* (how much rely on prior patents).
 - **Value**: (1) the number of times a patent is renewed by its owner and (2) *forward citations*.
 - *Hill, R., Yin, Y., Stein, C., Wang, X., Wang, D., & Jones, B. F. (2025). The pivot penalty in research. Nature, 642(8069), 999-1006.*
 - **Pivot size**: cosine distance between the references of the focal paper and prior papers.
 - **Impact**: (1) being top 5% cited within field-year and (2) citation count.

Backup: prior work

- Component familiarity
 - *Fleming, L. (2001). Recombinant uncertainty in technological search. Management science, 47(1), 117-132.*
 - **Component familiarity:** inventors' familiarity with components, based on the average degree to which the components have been recently and frequently used.
 - *Katila, R., & Ahuja, G. (2002). Something old, something new: A longitudinal study of search behavior and new product introduction. Academy of management journal, 45(6), 1183-1194.*
 - **Search depth:** the average number of times the firm has repeatedly used the citations.
 - **Search scope:** the proportion of previously unused citations.
- Combination atypicality
 - *Uzzi, B., Mukherjee, S., Stringer, M., & Jones, B. (2013). Atypical combinations and scientific impact. Science, 342(6157), 468-472.*
 - *Lin, Y., Evans, J. A., & Wu, L. (2022). New directions in science emerge from disconnection and discord. Journal of Informetrics, 16(1), 101234.*
 - 'The z-score of atypicality is deeply related to a common measure in information science, the **Pointwise mutual information (PMI)** between two items.'

Backup: data

- We analyze knowledge recombination in **scientific publications**.
 - Publication and citation data from SciSciNet-v2; journal vocabulary from OpenAlex.
 - Outcomes (disruption, field-year-normalized citation impact) precomputed in SciSciNet-v2.
- **Knowledge components:** journals cited by each paper.
 - Each paper is a set of cited journals: $p = \{k_1, k_2, \dots, k_n\}$; each cited journal k_i is a component.
 - **Frozen vocabulary** of $V = 28,360$ journals (OpenAlex is_core journals with > 100 works).
- **Researcher trajectories:** author-level publication histories ordered by date.
 - Prior history is the sequence of prior papers: $H_r = \{p_1, p_2, \dots, p_m\}$.
 - Only first-authored publications enter the trajectory, to capture the researcher's own path.
- **Sample restrictions:**
 - Journal articles (non-retracted) published **1970–2024**, team size ≤ 10 .
 - 10–100 references to indexed journals — excluding short or excessively long reference lists.
 - **Model-training:** researchers with ≥ 5 eligible papers (annual productivity less than ≤ 60).
 - **Analysis-time:** ≥ 5 th paper of each researchers (minimum prior history of 4 papers).

Backup: sample construction

Journal (vocabulary) filtering

Step	Sources
SciSciNet-v2 sources	260,811
Journals (<code>is_core</code> , >100 works)	28,360

Paper filtering

Step	Papers
Raw SciSciNet-v2 papers	249.8M
Journal articles, non-retracted	197.7M
Year 1970–2024, team ≤ 10	65.3M
10–100 indexed-journal refs	34.0M

Researcher filtering

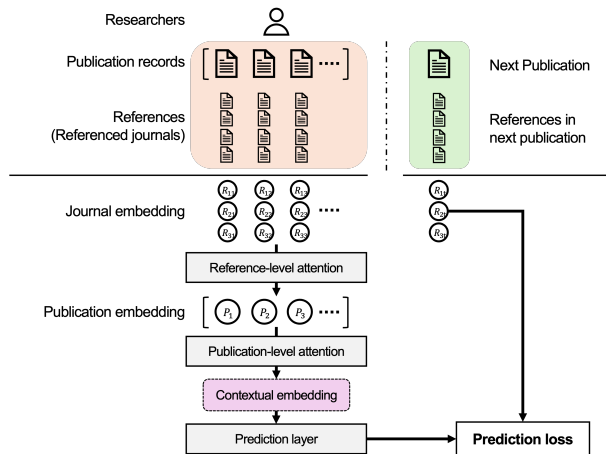
Step	Authors
With ≥ 1 eligible paper	20.0M
≥ 5 first-authored papers	2.0M

Annual eligible papers & researchers

Year	Papers	Researchers
1970	19,298	15,084
1980	62,764	47,551
1990	137,597	102,927
2000	302,119	214,377
2010	649,847	443,580
2015	849,619	570,289
2016	876,591	586,382
2017	895,578	597,600
2018	927,523	613,083
2019	966,482	628,774
2020	1,040,915	656,324
2021	1,052,069	653,095
2022	1,003,836	613,553
2023	914,295	563,392
2024	848,934	523,453

Total: $\sim 20.8M$ papers; 2.0M distinct researchers.

Backup: sequential model details



Hierarchical self-attentive model

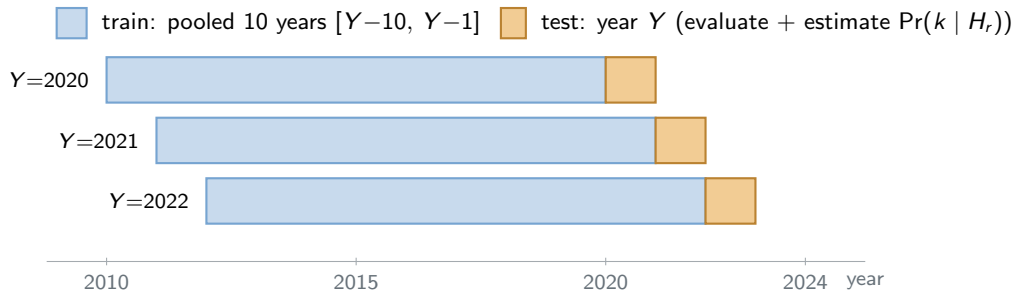
- Journal (reference) $e_k = \text{Emb}(k)$
- Paper $p_t = \text{Attn}(\{e_k\}_{k \in p_t})$
- Researcher $h_t = \text{Causal_Attn}(x_{1:t})$
- $\Pr(k | H_r) = \text{softmax}_k(\text{score}(h_t, e_k))$
- Weighted cross-entropy loss

$$\mathcal{L} = \sum_{k \in P_{t+1}} \frac{c_k}{\sum_c c_c} [-\log \Pr(k | H_r)]$$

Specification

- Journal embedding dim 2048; 1 layer; year + career-year embeddings.
- ≤ 64 recent papers for each researcher;
 ≤ 128 references for each paper.

Backup: data splitting for model training, evaluation, and estimation



- For each target year Y : pool the previous 10 years to train, then evaluate with first-authored papers published in Y . Window rolls across $Y \in [1980, 2024]$.
- For each focal paper, the input to the model is the researcher's prior history H_r (first-authored papers *strictly before* the focal paper).

Backup: the anchoring premium by career stage

	career stage		
	≤ 5 yr	5–15 yr	16–30 yr
Conventionality			
routine	+1.07***	+0.99***	+1.18***
pivot	+1.30***	+1.40***	+1.29***
leap	−0.23***	−0.21***	−0.33***
Share of atypical pairs			
routine	+0.88***	+0.89***	+1.21***
pivot	+0.48***	+0.52***	+0.47***
leap	+0.33***	+0.27***	+0.19***
observations	865,702	2,219,380	1,093,436
pseudo- R^2	0.143	0.128	0.134

Logit on top-5% impact; *** $p < 0.01$.

- The anchoring premium (**pivot** > **leap**) **grows with seniority**: $+0.16 \rightarrow +0.25 \rightarrow +0.28$.
- For **early-career** authors it is only marginal ($p = 0.04$); it sharpens with experience.

Backup: the anchoring premium holds by exploration distance (pivot size)

	reach distance (pivot size)		
	low ($\leq .3$)	mid (.3–.6)	high ($> .6$)
Conventionality			
routine	+1.26***	+0.96***	+0.77***
pivot	+0.99***	+1.57***	+1.97***
leap	−0.19***	−0.35***	−0.37***
Share of atypical pairs			
routine	+0.95***	+0.98***	+0.79***
pivot	+0.43***	+0.53***	+0.61***
leap	+0.22***	+0.25***	+0.40***
observations	1,940,300	1,729,854	735,644
pseudo- R^2	0.152	0.124	0.085

Logit on top-5% impact; *** $p < 0.01$.

- The anchoring premium (**pivot** > **leap** share) persist at **every distance**: anchoring pays even for the farthest reaches.

Backup: threshold-free test

- **Anchoredness of atypicality:** for each *atypical* pair in a paper, take the familiarity of its *more familiar* component, then average over all the paper's atypical pairs.

Predictor	coefficient
Conventionality	+2.52***
Share of atypical pairs	+1.34***
Anchoredness of atypicality	+0.16***
observations	5,851,709
pseudo- R^2	0.109

Logit on top-5% impact; *** $p < 0.01$.

- Anchoredness is **positive and significant**. The impact of papers increases when their atypical pairs are more anchored in familiar components.