

Categorical Engagement and the Contingent Nature of Typicality Effects

Alex Tyulyupo, Balázs Kovács

Yale University

Abstract: Market categories can favor typical members over atypical ones, yet this “categorical imperative” operates inconsistently across contexts. We argue that typicality effects depend on how evaluators engage with categories during search. Using behavioral simulation in which participants search a large database of companies for competitors, we distinguish between searches that explicitly invoke industry classifications and those using alternative methods, and measure goal-category congruence through semantic alignment between evaluators’ objectives and the categories they employ. We find that typical companies become more likely to be selected in category-based searches; other searches produce no typicality effects. Among category-based searches, goal-category congruence moderates the effect: typicality strongly predicts selection when goals and categories misalign but becomes irrelevant when they align well. These findings identify a specific microprocess through which categorical effects become contingent, offering a process-level explanation for variation documented across audiences, producer characteristics, and evaluation contexts.

Keywords: categories; search; entrepreneurship

Reproducibility Package: De-identified, model-ready data and code for reproducing all statistical tables and figures are available at (<https://doi.org/10.17605/OSF.IO/6Z8RS>). Raw Crunchbase records, company and venture descriptions, and participant-identifying information are not shared because of data-use and confidentiality restrictions.

Citation: Tyulyupo, Alex, and Kovács, Balázs. 2026. “Categorical Engagement and the Contingent Nature of Typicality Effects” *Sociological Science* 13: 884-914.

Received: May 6, 2026

Accepted: June 23, 2026

Published: July 28, 2026

Editor(s): Ari Adut, Clayton Chidress

DOI: 10.15195/v13.a34

Copyright: © 2026 The Author(s). This open-access article has been published under a Creative Commons Attribution License, which allows unrestricted use, distribution and reproduction, in any form, as long as the original author and source have been credited. 

MARKET categories give audiences a way to make sense of producers, sorting firms into recognizable types and supplying expectations about what each type should look like. A central concern in this tradition is how audiences treat producers that fit their category poorly, whether by spanning several categories at once or by sitting at the margins of a single one. The classic answer is that such atypicality carries a penalty. In Zuckerman’s (1999) categorical imperative, securities analysts overlooked firms whose industry membership was ambiguous, and those firms traded at a discount. Similar penalties have since been documented across many settings, from film producers who span genres (Hsu 2006) to restaurants with unconventional menus (Kovács and Johnson 2014). The mechanism is well understood: atypical offerings are harder to process because they do not match audiences’ mental representations of a category, and they unsettle the taken-for-granted order that classification systems carry (Douglas 1986).

However, the penalty is far from uniform, and its magnitude and even its direction vary widely across contexts. Venture capitalists (VCs) reward the categorical ambiguity that consumers penalize (Pontikes 2012), and high-status producers escape penalties that constrain lower-status ones (Zuckerman et al. 2003; Sgourev and Althuizen 2014). Accountability also changes how evaluators respond

to atypical organizations (Betancourt, Hoever, and Wezel 2025), and lenient or nascent classification systems generate no penalty for boundary spanning (Ruef and Patterson 2009; Pontikes and Barnett 2015). Synthesizing 25 years of category-spanning research, a recent meta-analysis finds that spanning has no context-independent effect on evaluations (Ahn, Vergne, and Sharkey 2026). The penalty for atypicality, in other words, is highly contingent.

Recent reviews organize this variation by identifying moderators rooted in audiences, producers, objects, and categorization systems (Cutolo and Ferriani 2024; Ahn et al. 2026). We complement these accounts by specifying one process-level mechanism generating such variation. Why does the same categorical structure produce penalties in one context and not in another? What practices of evaluation activate the typicality premium or allow actors to escape it?

We argue that answering these questions requires observing how actors engage with categorical structures during evaluation. Prior research relies primarily on archival data to measure outcomes like funding decisions or critical reviews, then infers that categories shaped these results. However, these methods cannot reveal whether evaluators actually consulted formal classifications, which specific categories they used, or whether their goals aligned with chosen categories. These engagement practices, we propose, determine when and how categorical structures produce the effects attributed to them.

This paper makes categorical engagement visible through behavioral simulation. We observe prospective entrepreneurs as they search a large database of company profiles to identify potential competitors for their venture concepts. Competitor identification is a strategically significant task. The competitors that an entrepreneur identifies shape subsequent assessments of market attractiveness, inform positioning decisions, and guide resource allocation (Porter 1980). This setting allows us to capture the practices through which categories come to matter: when entrepreneurs choose to search using industry classifications versus keywords, which specific industries they select, and how these choices shape selections of typical versus atypical firms. Unlike traditional outcome-only studies, our approach observes the process through which categorical structures are activated.

To analyze categorical engagement, we define three theory-level constructs. Categorical engagement is the practice of activating a formal categorical structure during search or evaluation. Goal-category congruence is the fit between evaluators' goals and the members of the categories they use. The focal category is the dominant category in the search results and the categorical frame relative to which typicality is assessed.

Our findings reveal that typicality effects depend on the search process itself. First, the typicality premium only appears in category-based searches; in searches that do not use industry categories as the primary semantic frame, typicality has no effect on selection. Second, among those who search by industry, venture-industry congruence transforms the typicality effect. When goals and categories misalign (low congruence), typical firms hold a selection advantage. When they align well (high congruence), this advantage disappears and atypical firms become equally likely to be selected. The same categorical structure produces different effects depending entirely on whether it is engaged and how well it fits the searcher's goals.

To validate this engagement mechanism, we examined whether typicality in nonfocal categories also predicts selection. Only typicality relative to the focal category affects selection, whereas typicality in other industries the firm belongs to shows no effect. This pattern confirms that typicality effects emerge from engaging with specific categorical structures.

In exploratory analyses, we also examine systematic variation in who achieves high congruence. Better-educated participants tend to generate results from more congruent industries. Conversely, holding pre-existing venture concepts appears to reduce congruence, a pattern consistent with cognitive entrenchment in mental models (Tripsas and Gavetti 2003). Although these patterns are only suggestive, they indicate that categorical navigation may be a capability that varies systematically across actors, with potential implications for who can identify nontraditional competitive threats or opportunities.

These findings advance research on market categories by demonstrating that typicality effects emerge through practices of categorical engagement.

Theory

The Typicality Premium and Its Moderators

Cognitive research shows that categories are organized around prototypes, with more typical members recognized faster (Rosch and Mervis 1975) and judged more favorably than atypical ones (Winkielman et al. 2006; Vogel et al. 2021). This prototype structure shapes evaluation in markets, where audiences often penalize firms that deviate from category norms, whether by spanning multiple categories (Zuckerman 1999; Hsu, Hannan, and Koçak 2009; Kovács and Hannan 2010, 2015) or appearing as atypical members of a single category (Hannan et al. 2019). The penalty is usually traced to cognitive friction, because atypical offerings are harder to process when they do not match mental category representations, and to the normative expectations that classification systems carry (Douglas 1986).

How strongly this premium operates, however, varies considerably. A growing body of research documents conditions under which it weakens or even reverses. The penalty for atypicality is contingent on the audience evaluating the actor. Pontikes (2012) finds that VCs reward what consumers penalize, a difference she attributes to their respective goals and expertise: VCs seek flexible opportunities and can parse novel combinations, whereas consumers prefer easily understood products. This variation extends beyond professional roles to audience cultural orientation. Goldberg, Hannan, and Kovács (2016) demonstrate that some cultural omnivores, whose status relies on mastering a variety of distinct categories, penalize atypical offerings that blur those same boundaries. Characteristics of the producer also matter. High-status actors can often cross category boundaries without penalty, as their established reputation provides a halo of legitimacy that compensates for deviance (Zuckerman et al. 2003; Sgourev and Althuizen 2014). Finally, the nature of the classification system itself is an important contingency. Penalties are strongest in mature, institutionalized category systems (Ruef and Patterson 2009), and a recent meta-analysis confirms that spanning penalties are stronger in bounded than in unbounded category systems (Ahn et al. 2026).

However, existing explanations for this variation focus primarily on characteristics of evaluators, producers, or contexts. Pontikes (2012) argues that VCs and consumers have different goals and expertise. A producer's career stage determines whether categorical conformity is rewarded; typecasting helps novices secure work but can later trap established veterans (Zuckerman et al. 2003). Market maturity determines whether boundaries are enforced (Ruef and Patterson 2009). These studies identify important moderators, but they rely on archival data measuring evaluation outcomes and work backward to infer that categories shaped these results. The practices through which evaluators engage with categorical structures before reaching these outcomes remain unobserved.

This gap matters because categorical engagement may not be constant, and this variability might explain when categorical effects emerge. Kim and Jensen (2011) show that opera companies can appeal both to audiences who value typicality and to audiences who appreciate atypicality by strategically ordering their catalogs. More broadly, the documented moderators hint that different evaluators may activate and use categories differently. Pontikes's finding that VCs appreciate category-blurring could reflect not just different preferences but different practices. VCs conducting extensive due diligence may explore multiple categorical frames, while consumers rely on platform interfaces that make certain categories salient. Audiences might be less aware of nascent categories (Ruef and Patterson 2009), might not use them for search, and therefore remain oblivious to companies' fit in their industries. Without observing how actors engage with categories, we cannot distinguish whether variation in categorical effects stems from inherent differences in preferences or from differences in engagement practices.

The Search Function of Categories and Its Alternatives

Observing these engagement practices means attending to how actors search. Identifying competitors is a sensemaking task whose process is still not well understood (Gur and Greckhamer 2019), one in which the relevant rivals are constructed rather than read off the market (Porac and Thomas 1990; Cattani et al., 2018; Arora-Jonsson et al. 2021; Sands et al. 2021). Established managers build their sense of competition from inputs that tend to make their assessments converge across similar firms, including data on consumer substitution, resource-based assessments of rivals, and the shared mental models that emerge through direct interaction among producers (Chen 1996; Clark and Montgomery 1999; Porac, Thomas, and Baden-Fuller 2011). Nascent entrepreneurs lack all of these and must rely instead on public cognitive artifacts, of which industry classifications are the most visible.

Entrepreneurs identifying competitors, like managers searching for acquisition targets or investors searching for opportunities, must decide how to organize their search over a large and complex option space. Before digitization, actors relying on formal information systems such as library stacks, industry directories, and patent classifications had to navigate established categorization systems to find anything. Categories were constant features of the information environment. Today, however, digitization has made categorical engagement variable. Actors can search using formal categorical structures, where industry classifications, occupational

categories, and product types serve as filters that reduce vast possibilities to manageable consideration sets (Kovács, Carnabuci, and Wezel 2021). However, they can also search using keywords, follow social recommendations, examine popularity rankings, or combine multiple approaches.

The choice of a search path has consequences beyond determining which items appear in results. It shapes the cognitive frame through which those items get evaluated. We propose that categorical engagement determines whether typicality effects emerge through a cognitive priming mechanism. When actors choose to search using formal categories, they activate that classification system as an organizing lens. Cognitive research shows that explicit categorical framing changes which features become salient in similarity judgments (Tversky 1977; Goldstone 1994). By searching within a category, the actor primes attention to category-defining features, making typicality within that industry a relevant evaluative criterion.

This priming operates regardless of awareness. When an entrepreneur searches within “Education,” they implicitly accept that industry’s organizing logic, in which firms are meaningfully compared by how well they match educational prototypes like universities or online course platforms. In contrast, when actors search using alternative methods such as keyword matching for “interactive learning,” location filters, or following network recommendations, they do not invoke formal categorical structures as organizing lenses. These searches may surface similar firms, but the evaluative frame differs. A keyword search primes attention to the specific features expressed in the query. In such a search, the entrepreneur pursues the goal-based category “competitors to my venture” without reliance on ready-made categories (Barsalou 1983). Evaluation focuses on whether firms offer the searched-for attributes, which may or may not correlate with typicality in any particular industry classification.

Consider how patent examiners search for prior art. When they search using classification codes, they retrieve results that are categorically focused, and their evaluation should attend to how well an invention fits within that class (Kovács et al. 2021). When they search by keywords describing technical features, they may find similar patents without the same categorical framing. Similarly, a restaurant reviewer using a platform organized by cuisine categories (“French Cuisine” vs. “Chinese Cuisine”) becomes attuned to how well restaurants exemplify their assigned category, whereas one discovering restaurants through a friend’s recommendations may not consider typicality at all. The act of engaging a categorical structure shapes what becomes the basis for evaluation.

This perspective departs from traditional views that treat categories as an exogenous structure determining evaluative outcomes. We propose instead that a formal category is one evaluative resource among several, and that an actor’s search approach governs how much weight it carries. This complements work showing that audiences evaluate not only through prototype-based logic but also through goal-based logic, organized around the actor’s goal rather than fit to a prototype (Barsalou 1983; Durand and Paoletta 2013; Glaser, Atkinson, and Fiss 2020). Searching within a formal category makes prototype-based evaluation salient, whereas the choice to search by other means indicates that framings other than the category appeared more appropriate, leaving comparison to be organized by the goal itself.

For producers, understanding that different search pathways activate different evaluative frames points to opportunities to position their offerings within beneficial categorical framings.

When Typicality Effects Emerge: The Role of Categorical Engagement

This framework leads to our first hypothesis about when typicality effects emerge. Prior research has documented that audiences penalize atypical category members, but these studies measure evaluation outcomes without observing whether formal categories were engaged during the selection process. If typicality effects require categorical engagement, then they should only appear when actors explicitly use categories to structure their search. When actors use alternative search methods that do not invoke formal categorical structures, the dimensions defining typicality should not be primed, and typicality should not systematically affect selection.

The mechanism is straightforward: searching within a category makes that category's prototype the implicit reference point. An entrepreneur who filters for "Education" companies is prompted to evaluate how well each firm represents an Education company. Highly typical firms, that is, those that clearly exemplify the category prototype, enjoy an advantage because they fit the activated evaluative frame. Atypical firms, even if substantively relevant, deviate from the prototype and thus fare worse under this categorical lens. However, an entrepreneur who searches using keywords like "interactive learning platform" or "professional certification" does not activate the Education industry as a reference point. The same firms may appear in results, but evaluation centers on whether they offer interactive learning or certifications rather than whether they exemplify the Education industry prototype.

The choice of search path can have two sources. The entrepreneur may hold the Education category but search by other means, or may lack an industry concept that maps well enough onto the venture to use as a filter. Searching with a category presupposes some familiarity with it, so the absence of category-based search may itself signal that no fitting industry concept was available. Our mechanism predicts the same outcome in either case. No industry prototype is activated as the evaluative lens, so typicality within any particular classification becomes irrelevant.

Hypothesis 1: When evaluators search using industry categories, more typical firms will be more likely to be selected; this typicality effect will not appear in searches that bypass formal categories.

How Typicality Effects Operate: The Moderating Role of Goal-Category Fit

Engaging with categorical structures is necessary for typicality effects to emerge, but it is not sufficient to determine how strongly those effects operate. We argue that the fit between an actor's goal and their chosen category determines the magnitude of typicality effects. This argument recognizes that categorical engagement is purposeful and actors use categories to accomplish goals like finding relevant

competitors, identifying opportunities, or evaluating options. Whether a category serves these goals effectively depends on alignment between what the actor seeks and what the category contains.

Goal-category congruence captures the semantic alignment between the actor's goal and the members of the searched category. In the case of early competitor identification, high congruence means the entrepreneur has found an industry whose members are semantically similar to their venture. Low congruence means they are searching in a poorly fitting category where most firms are semantically distant from what the entrepreneur seeks. Similarly, for consumer or expert audiences, a high-congruence category is one that contains many items that could satisfy the immediate goal of the search: a genre or cuisine that aligns with the consumer's cultural preferences, an industry that aligns with the expertise of a security analyst, and so on.

Congruence shapes whether evaluation operates at the category level or the item level. When congruence is low—whether the entrepreneur deliberately chose a broad category for strategic reasons or simply failed to identify a more fitting one—selection focuses on categorical representation, identifying which firms best exemplify what this category of competition offers. When congruence is high, entrepreneurs can move beyond categorical representation and focus on comparing individual competitors on their specific merits.

Consider an entrepreneur developing an all-in-one workspace platform that integrates documents, databases, and project management. If this entrepreneur searches in “Productivity Tools,” they face low congruence. While their platform relates to productivity, the productivity software industry is semantically centered on traditional office suites, such as word processors, spreadsheets, and presentation tools. In this low-congruence context, typicality becomes decisive. The entrepreneur selects firms that best exemplify the productivity software category—comprehensive office suites that represent the categorical prototype. A platform that blends wikis, databases, and flexible workspaces is atypical for productivity software and gets screened out not because it is irrelevant, but because it fails to represent what “productivity software” as a category means. The evaluation operates at the category level: which firm best exemplifies this domain?

The evaluation context transforms when the same entrepreneur searches in “Collaboration Tools,” achieving high congruence. This category is semantically centered on flexible, integrated platforms for team coordination, closely aligned with the entrepreneur's venture. Now, the encountered firms are viewed as individual direct competitors rather than representatives of an adjacent competitive category. A platform with an unconventional approach to workspace organization simply represents a strategic variant within a relevant competitive space. Typicality within the category becomes irrelevant to selection.

The distinction between category-level and item-level evaluation explains why typicality effects emerge in low-congruence contexts but not in high-congruence contexts. When searching in a poorly fitting category, entrepreneurs are trying to capture what that category of competition represents, making typicality the relevant heuristic for identifying the best exemplar. When searching in a well-fitting category, entrepreneurs are comparing specific competitors, making individual attributes more relevant than categorical typicality.

Hypothesis 2: Among searches that engage industry categories, high goal-category congruence will neutralize the typicality effect.

This mechanism is not unique to entrepreneurial competitor search. Consider a consumer planning weekend activities who prefers bars. When this person searches among “museums” for cultural options they might enjoy, they face low congruence, since museums are semantically centered on historical artifacts, art collections, and educational exhibits. The categorical frame of “museums” does not meaningfully capture what this person seeks. In this context, they are likely to select the most typical museum in town. An interactive gin-tasting space with museum-like educational elements would be atypical for the museum category and likely overlooked, not because it is irrelevant to their interests, but because it fails to exemplify what “museums” as a category represents.

However, if the same person searches among “bars,” they achieve high congruence, since bars align closely with their preferred activities. That same gin-tasting establishment, which is also a member of the “bar” category, gets evaluated on its individual merits. Notably, the establishment is as atypical of a bar as it is of a museum, yet this atypicality becomes irrelevant under high congruence. The establishment has not changed; what has changed is the categorical lens, which determines whether typicality matters for selection.

Together, these hypotheses suggest that what prior research has treated as structural properties of categories can emerge through practices of engagement. The same classification system produces different effects depending on whether actors engage it and, when they do engage it, how well their goals align with the categories they choose. This reframes categorical constraints not as inherent properties but as consequences of how actors navigate categorical structures during evaluation.

Empirical Setting and Sample

Behavioral Simulation Methodology

To test our theory of how categorical engagement shapes typicality effects, we conducted a behavioral simulation (Dutton and Stumpf 1991), observing participants’ search and selection patterns as they identified competitors from over 500,000 companies on our custom-built platform. We use the same platform and behavioral-simulation design that we introduced in a companion study examining how geographic categories and location-typicality shape competitor identification (Kovács and Tyulyupo 2026), though the present study draws on a distinct sample¹ and focuses on industry categorization. As McCall and Lombardo (1982:533) note, the primary goal of behavioral simulation is to “mirror the environment [while] allowing free behavior within it.” Behavioral simulations have been used to study complex strategic questions in management research, including competitive strategy (Clark and Montgomery 1999) and entrepreneurship-specific problems (Katila, Chen, and Piezunka 2012).

Our methodology preserves the voluntary nature of categorical choices: participants can choose whether to search using industry categories or other methods,

select which specific industries to employ, and determine how long to search. Simultaneously, we maintain knowledge of each participant's specific goal (their venture concept), enabling us to measure venture-industry congruence, the fit between an actor's goal and the categorical tool they employ.

Alternative methodologies would not provide the observational precision our theory requires. Traditional experiments pre-assign participants to categories, thereby eliminating the voluntary engagement choices central to our hypotheses. Observational data from commercial platforms, while preserving natural search behavior, would lack information about searchers' underlying goals, making it impossible to assess goal-category congruence. The behavioral simulation thus combines ecological validity with measurement precision, allowing us to observe both the practice of categorical engagement and its consequences for evaluation.

We chose entrepreneurial competitor search as the empirical setting because it creates ideal conditions for observing categorical engagement and its consequences. Accurately identifying competitors is consequential, as the competitors an entrepreneur selects become reference points shaping market attractiveness assessments, positioning decisions, and resource allocation (Porter 1980; Porac et al. 2011; Clark 2011). Yet this task is cognitively demanding (Gur and Greckhamer 2019), particularly for nascent ventures where the field of potential competitors is effectively borderless, with hundreds of thousands of possible competitors. The number of options far exceeds cognitive capacity, requiring simplifying heuristics to structure the search (Simon 1955).

Consistent with this account, nascent entrepreneurs face this task as outsiders and lean on public categorical structures rather than the insider information available to established firms. This reliance makes their categorical engagement observable and, as with other audiences who evaluate market actors through such structures, creates conditions where typicality effects can emerge.

This setting also provides observational advantages. Because participants define their own venture concepts, we can precisely measure the fit between their goals and the categorical tools they employ. This goal clarity allows us to operationalize venture-industry congruence and observe how goal-category congruence transforms the evaluation of atypical firms.

Sample

We recruited 668 participants through the behavioral laboratory of the business school at a private U.S. university. The study was advertised as "A study about entrepreneurship," attracting individuals with interest in and exposure to entrepreneurial thinking. Institutional Review Board approval was obtained prior to data collection. The study was conducted on computers in a supervised laboratory setting, which helped reduce potential distractions and ensure adherence to the protocol. Participants received \$20 for their time. The search portion of the study took eight minutes on average. After excluding respondents who reported using external information sources or who engaged with the search task for less than three minutes, 567 participants remained in the final sample.

This sample is appropriate for testing our theory because the cognitive mechanisms we study operate independently of whether participants intend to launch

these ventures. Demographically, 57 percent of participants were female, and the average age was 29. Education levels varied: 43 percent held bachelor's degrees, 33 percent had completed graduate degrees (master's or doctoral), 13 percent had some college or associate's degrees, and 7 percent had high school diplomas.

The Competitor Identification Task

The study began by asking each participant to define their entrepreneurial goal. They were prompted to propose a venture concept and provide a short description of it. Participants could either select one of six predefined prompts or propose an original idea, as 38 percent of participants chose to do. The predefined prompts were: (1) cloud-based data analytics platform, (2) sustainable urban development firm, (3) fashion store, (4) solar panel installation and maintenance, (5) online vocational training courses, and (6) wine store.

In all cases, participants were required to write a description of the venture (minimum 100 characters), explaining "what product or service it offers." The requirement for individualized descriptions ensured variation even among participants selecting the same predefined prompt. For instance, two participants selecting the "wine store" prompt might describe it differently—one as a local shop focused on education and the other as a subscription-based e-commerce platform. This specificity is central to our design, as what constitutes a high-congruence industry depends entirely on the unique description of the venture.

After describing their ideas, participants were tasked with searching for potential competitors using a research platform we developed for this study. The platform accessed real company profiles sourced from Crunchbase under an educational-research agreement with our university. Its interface resembled Crunchbase, offering company descriptions and other core details.

Participants could search for competitors using four filters: keyword, industry, number of employees, and location. These filters could be used individually or in combination. For example, a participant could search using only the keyword "wine tasting," only the industry category "Food and Beverage," or both simultaneously along with a location filter. A search returned up to 20 company profiles per page from the database and presented them in randomized order. Participants could paginate through results, refine their queries by adjusting filters, inspect individual company profiles by clicking on them for detailed information, and build a final list of competitors by selecting companies they deemed relevant.

Figure 1 shows an example search interface and results for a participant who proposed the idea of a local wine store with a strong educational element. The interface displayed the participant's venture description at the top to maintain focus on their goal, showed the current search parameters, and listed resulting companies with their descriptions and key attributes. Participants were encouraged to identify approximately six competitors but could select more or fewer based on what they found relevant. While a timer was visible to help participants pace themselves, no strict time limit was enforced. This design allowed participants to search freely while maintaining reasonable task boundaries.

The categorical structure participants navigated was Crunchbase's proprietary classification system, which contains 747 distinct industry categories. Our

Company Insight Hub

All Companies My selected Companies Settings Logout

Company Lists

Your task is to find and rank **about 6** potential competitors for:

An online platform offering training courses for high-paying careers, including pre-recorded videos and personalized support.

Search for **about 8 minutes**

Please wrap up the search soon.

Description Keywords
E.g. Natural Gas, R4 Capi

Number of Employees
11 - 50

Headquarters Location
E.g. Japan, Boston, Europ

Industry
Winery

Search

Reset

I'm done with selecting the companies

Organization Name	Industries	Headquarters Location	Description	Employee size
Bellview Winery	Manufacturing,Wine And Spirits,Winery	United States of America, New Jersey, Landisvill...	Bellview Winery manufactures and distributes wines.	11-50
Tablas Creek Vineyard	Farming,Wine And Spirits,Winery	United States of America, California, Paso Roble...	Tablas Creek Vineyard provides offers a rhone varietal budwood to growers.	11-50
Bucher Vaslin North America	Wine And Spirits,Winery	United States of America, California, Santa Rosa	Bucher Vaslin North America is a winemaking equipment company.	11-50
Phase 2 Cellars	Food and Beverage,Wine And Spirits,Winery	United States of America, California, San Luis Obispo	Phase 2 Cellars is a producer of nuanced and balanced pinot noir and chardonnay wines. Founded in 2007.	11-50

Figure 1: Information portal interface with example search results.

theoretical interest centers on whether and how participants choose to engage this categorical structure. When participants use the industry filter, they explicitly invoke formal categories to organize their search, activating industry classifications as an evaluative lens. This categorical engagement is observable and voluntary: participants can choose to search using industry categories, use other methods like keywords, or combine approaches. By observing these choices, we can test whether typicality effects emerge only when actors engage categorical structures and whether the fit between chosen categories and searcher goals transforms how typicality operates.

Typicality Effects and Categorical Engagement

Model and Dependent Variable

To test Hypotheses 1 and 2, we employed logistic regression models with participant fixed effects. The unit of analysis is each instance where a company appeared in a participant's search results ($n = 76,814$). This modeling approach allows us to examine how various factors influence the likelihood that a participant selects a given company as a competitor, while controlling for unobserved heterogeneity among participants. By comparing choices within the same participant across different search contexts, we can test that typicality effects are not uniform but depend on whether actors engage categorical structures and on the fit between their goals and focal categories. Table 1 presents descriptive statistics and pairwise correlations of the variables included in the models.

Our dependent variable is *Company Selection as a Competitor*. It captures whether a participant identified a specific company as a competitor to their venture concept during the selection process, given that the company appeared in search results. The selection is recorded as a binary outcome for each company's appearance: a value

Table 1: Descriptive Statistics and Correlations.

Variable	Mean	SD	Min	Max	Median	1	2	3	4	5	6	7	8	9	10
1. Selection of company as competitor	0.03		0	1		1	0.05	-0.03	0.02	0.02	0.11	0.03	0.01	0.01	0.03
2. Company typicality	51.95	29.1	1	100	57.55	0.05	1	-0.16	-0.04	0.06	0.27	-0.01	0.05	-0.01	0
3. Category-based search	0.54		0	1		-0.03	-0.16	1	0.27	0.07	-0.21	0.03	0.02	0.06	0.08
4. Venture-industry congruence	0.26	0.08	-0.03	0.45	0.26	0.02	-0.04	0.27	1	0.34	0.44	0.05	-0.03	0.01	0.05
5. Venture typicality	56.37	23.39	1	90.5	62.4	0.02	0.06	0.07	0.34	1	0.18	0.01	-0.01	0.05	0.03
6. Venture-company semantic similarity	0.31	0.13	-0.13	0.89	0.3	0.11	0.27	-0.21	0.44	0.18	1	0.02	-0.05	-0.04	0
7. Company in United States	0.61		0	1		0.03	-0.01	0.03	0.05	0.01	0.02	1	-0.05	0.1	0.35
8. Number of employees	358.86	1633.65	5	15000	30	0.01	0.05	0.02	-0.03	-0.01	-0.05	-0.05	1	0.05	-0.02
9. Employee count filter used	0.21		0	1		0.01	-0.01	0.06	0.01	0.05	-0.04	0.1	0.05	1	0.22
10. Location filter used	0.28		0	1		0.03	0	0.08	0.05	0.03	0	0.35	-0.02	0.22	1

Note: Standard deviations and medians are omitted for binary variables.

of 1 indicates that the participant selected the company as a competitor, while 0 indicates they did not. On average, participants selected 6 percent of the companies appearing in their search results; at the observation level, the selection rate was 3 percent.

Venture-Industry Congruence

Venture-industry congruence operationalizes goal-category congruence in this study, capturing the semantic fit between the participant's venture description and the industry category that dominates their search results.

On each search results page, the focal industry is the one to which the largest number of returned firms belong. We calculate congruence relative to this focal industry. Because participants typically generated multiple search pages, the focal industry varies within a participant across pages. For example, if a search returns 20 companies, of which 18 belong to "E-Commerce," 10 to "Food and Beverage," and 3 to "Retail," E-Commerce is the focal industry for that page. Because companies can belong to multiple industries, these counts need not sum to 20. In the 10 percent of cases where multiple industries were equally frequent (usually semantically related ones), we determined the focal industry to be the one listed first in the first company's profile, maintaining measurement consistency across category-based and keyword-based searches. For each search results page, congruence is the mean semantic similarity between the participant's venture description and the descriptions of every company in our database that belongs to the focal industry, not only the firms appearing on that page. Using the mean captures the overall relevance of the category: a high-congruence industry is one where members are, on average, a good semantic match, while a low-congruence industry is noisier and contains more irrelevant firms.

To implement this measure, we generated dense vector representations (embeddings) for each venture description and each company description using the all-MiniLM-L6-v2 sentence transformer model, which is optimized for semantic similarity and clustering (Reimers and Gurevych 2019). This model maps text into a 384-dimensional space that captures semantic meaning. We then computed the cosine similarity between the venture's embedding and the embedding of each company in the focal industry across our entire database, and averaged these similarities to produce our venture-industry congruence score.

Venture-industry congruence is an industry-level construct measuring the fit between a searcher's goal and a categorical grouping. This measure applies regardless of whether participants explicitly used the industry filter in their search. The focal industry emerges from the composition of search results, whether generated by keyword search, industry filter, or a combination of approaches. To isolate this industry-level effect, we also calculate the direct semantic similarity between each venture description and each individual company that appears in the search results. This company-level similarity is included as a control in all models, ensuring that congruence effects reflect the overall fit of the predominant category rather than just individual company matches. Table 2 provides illustrative examples, showing high-congruence and low-congruence industry matches for several venture concepts from our study.

Table 2: Examples of Least- and Most-Congruent Venture-Industry Pairs.

Venture Concept	Least Congruent Industry	Score	Most Congruent Industry	Score
An autonomous hull-cleaning device for large ships, including cargo ships, tankers, cruise lines, and military vessels. The device would be able to clean while the ship is in motion, as opposed to only during docking.	Manufacturing	0.12	Marine transportation	0.28
The venture would be a wine store with half the store offering bottles for purchase and half serving as a seating area where folks can come in for happy hours and tastings. There would be cheese, crackers, and other small apps served, along with a lounge area for relaxation.	Service industry	0.12	Winery	0.34
These pickleballs resist cracking and misshaping in cold and warm weather; these are pickleballs that do not break from wear and tear. The purchase of only a couple of balls should last a lifetime, not three weeks like the current ones being used.	Dental	−0.03	Sports	0.04
My venture would identify abandoned and/or at-risk historic houses, purchase and rehabilitate them, and ultimately resell them to new owners who will preserve their legacy.	Health care	0.06	Home renovation	0.24
Reselling/renting furniture in high-density university towns. Offers a more environmentally friendly option for students moving to the area for 1–2 years. Reduces waste in the area. Can be more affordable than buying new IKEA/target or more expensive products.	Biotechnology	0.04	Furniture	0.27

Company Typicality

Company typicality measures how representative a company is of its industry category. We estimated the typicality of each company based on the industry most frequently represented on the search results page where the company appeared—the same focal industry used to calculate Venture-Industry Congruence. By aligning typicality with the dominant industry per search page, we capture the category most salient to participants during their search and selection process. Prior research often lacked precise knowledge of which category the audience was focusing on, using averages of typicality across multiple categories (Kovács and Johnson 2014), the number of categorical memberships as a proxy for overall typicality (Hsu et al. 2009), or simplified categorization systems without multiple category memberships

(Le Mens et al. 2023b). Our method addresses this by measuring typicality within the category that actually structured participants' evaluative context.

To compute typicality scores, we used the ChatGPT API (model GPT-4o) to rate how typical each company's description is for the relevant focal industry. Our approach builds on recent validation evidence showing that large language models produce accurate estimations of human typicality perceptions (Le Mens et al. 2023a), a method highlighted as an important innovation in the sociology of entrepreneurship (Botelho, Gulati, and Sorenson 2024). Le Mens et al. (2023a) demonstrated that GPT-4 outperforms previous approaches to typicality estimation, including BERT-based embedding methods (Le Mens et al. 2023b), without requiring researchers to pretrain models on potentially biased datasets. Following their procedure, we issued the following prompt for each company-industry pair:

How typical is this company for the {industry label} industry? Provide the answer as a number from 1 to 100, where 1 is very atypical, and 100 is very typical. Drop any explanations, provide just the number. Your answer should contain no characters, ONLY NUMBERS. Use only this provided text description to make the estimation. The description of the company is {company description}

We modified the prompt by adding the instruction to “use only this provided text description to make the estimation” to avoid biasing typicality scores with additional information about specific companies that GPT-4o might possess but that was unavailable to our participants. Table A1 in the online supplement provides examples of twenty randomly selected company descriptions from the Education industry along with their estimated typicality scores, offering qualitative validation of the measure. The most typical company is a classic university, companies of average typicality focus on niche audiences or have unconventional organizational models (e.g., a cooperative network), and the least typical companies are research institutions that offer education as a secondary service.

In total, we calculated typicality scores for approximately 32,000 company-industry pairs, covering about 29,000 unique company profiles that appeared in participants' search results. These profiles constitute the risk set of companies that participants could have selected. This represents an important methodological advantage: most studies of typicality effects do not know which firms entered the consideration set and must assume that all category members were eligible for evaluation. This assumption is unrealistic, since most categories contain thousands of members and evaluators necessarily filter to a smaller set. By observing actual search results, we measure typicality only for companies that participants encountered, providing a more precise test of how typicality shapes selection within the relevant choice context.

To validate this measurement approach for our context, we conducted a human rating study comparing GPT judgments to human typicality perceptions. We randomly sampled 100 companies from our dataset using stratified sampling across the full range of GPT-estimated typicality scores. We recruited 220 U.S.-based participants with undergraduate degrees or higher via Prolific, each of whom rated the typicality of 10 randomly assigned companies on the same 1-100 scale. After removing seven participants who failed an attention check, all 100 companies in our

validation sample received at least 10 human ratings. We then assessed reliability by comparing GPT estimates to human ratings. GPT's correlation with averaged human ratings was $r = 0.72$ (95 percent confidence interval [0.61–0.80]). This result supports the validity of GPT-based typicality measurement for our analysis.

Category-Based Search

Category-based search is a binary variable indicating whether participants used the industry filter without the keyword filter to generate a particular set of search results. This operationalization isolates searches in which the industry category provides the primary semantic frame, rather than sharing that role with a keyword query. Participants may have used other filters, such as employee count or location, in combination with the industry filter. This distinction is foundational to our analytical strategy. Hypothesis 1 posits that typicality effects emerge only when actors explicitly engage formal industry classifications during search. In these searches, participants invoke categorical structures as an organizing lens, making industry-based typicality a relevant evaluative criterion. When they search by keywords, we theorize that typicality within industry classifications should not systematically affect selection.

Hypothesis 2 further proposes that among category-based searches, the strength of typicality effects depends on venture-industry congruence, that is, the fit between the searcher's goal and the chosen category. Consequently, our main analyses testing H2 focus on the subset of observations from category-based searches ($n = 41,211$), where categorical engagement has occurred and congruence can meaningfully moderate typicality effects.

Control Variables

To isolate the effects of typicality, congruence, and category focus, our models include participant fixed effects and several control variables. The participant fixed effects absorb all time-invariant characteristics, including venture concepts, general preferences, search duration, and any other stable individual traits, ensuring our coefficients reflect within-person variation in selection behavior across different search contexts. The control variables account for characteristics of the companies, the search context, and the search process itself.

First, we control for venture-level and company-level attributes. *Venture typicality* within the focal industry is calculated using the same GPT-4o method as company typicality but applied to the participant's own venture description. This accounts for the possibility that the typicality of the venture concept itself influences selection patterns. *Venture-company semantic similarity* captures the direct textual fit between each participant's venture description and each individual company's description, calculated using the embedding approach described earlier. This company-level control ensures that the effects of industry-level congruence and typicality are not simply artifacts of dyadic similarity between specific venture-company pairs. We also control for basic company characteristics: *Company location* includes a binary indicator for whether the company is located in the United States, and *Company size* is measured by the number of employees.

Second, we control for factors related to the search interface and process. *Position in search results* captures the company's rank (1–20) on the results page; prior research shows that ranking heavily influences selection (Ghose, Ipeiritis, and Li 2014). *Number of results on page* controls for how many companies appeared on that search page, as page density may affect attention and selection rates. *Search filter usage* includes binary indicators for whether the participant used the employee count filter or location filter in generating that set of results, accounting for different search strategies that may systematically affect which companies appear and how they are evaluated.

Results

We now test whether typicality effects depend on categorical engagement (H1) and goal-category congruence (H2). Table 3 presents the results. Model 1 establishes a baseline by examining the average effect of typicality on selection without accounting for how participants searched. Mirroring prior research that has not observed search practices, this model treats typicality as a uniform predictor. We find a positive and significant coefficient for company typicality ($\beta = 0.006, p < 0.01$), indicating that more typical firms are more likely to be selected in a model that does not account for whether or how participants engaged categorical structures.

Model 2 adds category-based search and venture-industry congruence as main effects, along with controls for search parameters. This model allows us to assess whether engaging industry categories and achieving goal-category congruence independently affect selection, setting up the interaction tests. The coefficient for company typicality remains positive and significant ($\beta = 0.006, p < 0.01$), whereas venture-company semantic similarity shows a strong effect ($\beta = 5.674, p < 0.01$), confirming that direct dyadic relevance heavily influences which companies participants select as competitors.

Model 3 tests Hypothesis 1 by introducing the interaction between company typicality and category-based search. The results provide clear support for our proposition, as the main effect of company typicality becomes nonsignificant, while the interaction term is positive and significant ($\beta = 0.008, p < 0.01$). This pattern demonstrates that typicality effects emerge specifically in category-based searches. When participants search by keywords or other methods without engaging formal industry classifications, company typicality does not systematically affect selection. The effect of venture-company semantic similarity remains strong ($\beta = 5.567, p < 0.01$), showing that the categorical engagement effect operates beyond dyadic relevance. These results confirm that typicality premiums are not ambient features of market structure but emerge through the practice of categorical engagement, providing strong empirical support for focusing our congruence analysis (H2) on the subset of category-based searches where categorical structures have been activated.

Model 4 includes only category-based searches ($n = 41,211$) and introduces the interaction between company typicality and venture-industry congruence to test H2. The main effect of company typicality in this subsample is positive and significant ($\beta = 0.027, p < 0.01$), confirming the pattern observed in Model 3: within category-based searches, typical firms enjoy selection advantages. However, this effect is qualified by a significant negative interaction with venture-industry congruence

Table 3: Consequences of Venture-Industry Congruence.

	Dependent Variable: Competitor Selected			
	(1)	(2)	(3)	(4)
Company typicality	0.006 [†] (0.001)	0.006 [†] (0.001)	0.002 (0.002)	0.027 [†] (0.005)
Category-based search		−0.200 (0.147)	−0.678 [†] (0.183)	
Venture-industry congruence		−0.462 (1.071)	−0.313 (1.056)	2.961 (2.275)
Venture typicality	0.008* (0.004)	0.009* (0.004)	0.009* (0.004)	−0.003 (0.006)
Venture-company semantic similarity	5.654 [†] (0.295)	5.674 [†] (0.298)	5.567 [†] (0.302)	5.985 [†] (0.423)
Company in United States	0.347 [†] (0.072)	0.278 [†] (0.075)	0.281 [†] (0.075)	0.157 (0.113)
Number of employees (1,000s)	0.040* (0.020)	0.040* (0.020)	0.040* (0.020)	0.040 (0.030)
Position in search results	−0.061 [†] (0.005)	−0.061 [†] (0.005)	−0.061 [†] (0.005)	−0.048 [†] (0.007)
Number of results on page	−0.078 [†] (0.009)	−0.065 [†] (0.010)	−0.064 [†] (0.010)	−0.072 [†] (0.017)
Employee count filter used		0.076 (0.193)	0.075 (0.191)	0.154 (0.310)
Location filter used		0.736 [†] (0.159)	0.740 [†] (0.160)	0.842 [†] (0.253)
Company typicality × category-based search			0.008 [†] (0.002)	
Company typicality × venture-industry congruence				−0.061 [†] (0.018)
Constant	−4.597 [†] (0.486)	−4.968 [†] (0.553)	−4.686 [†] (0.541)	−5.431 [†] (0.834)
Participant FE:	Yes	Yes	Yes	Yes
Observations	76,814	76,814	76,814	41,211
Log likelihood	−9,855.933	−9,830.888	−9,819.657	−4,700.510
Akaike inf. crit.	20,859.870	20,817.780	20,797.310	10,151.020

Note: [†] $p < 0.01$; * $p < 0.05$.

Standard errors clustered by participant.

($\beta = -0.061, p < 0.01$). This interaction supports H2, demonstrating that the strength of typicality effects depends on goal-category congruence.

Figure 2 plots the predicted probabilities of selection from the interaction in Model 4 using the category-based search subsample, showing the probability of selection across levels of venture-industry congruence at three illustrative levels of company typicality (mean and ± 1 SD), with all other variables held at their mean or reference values.

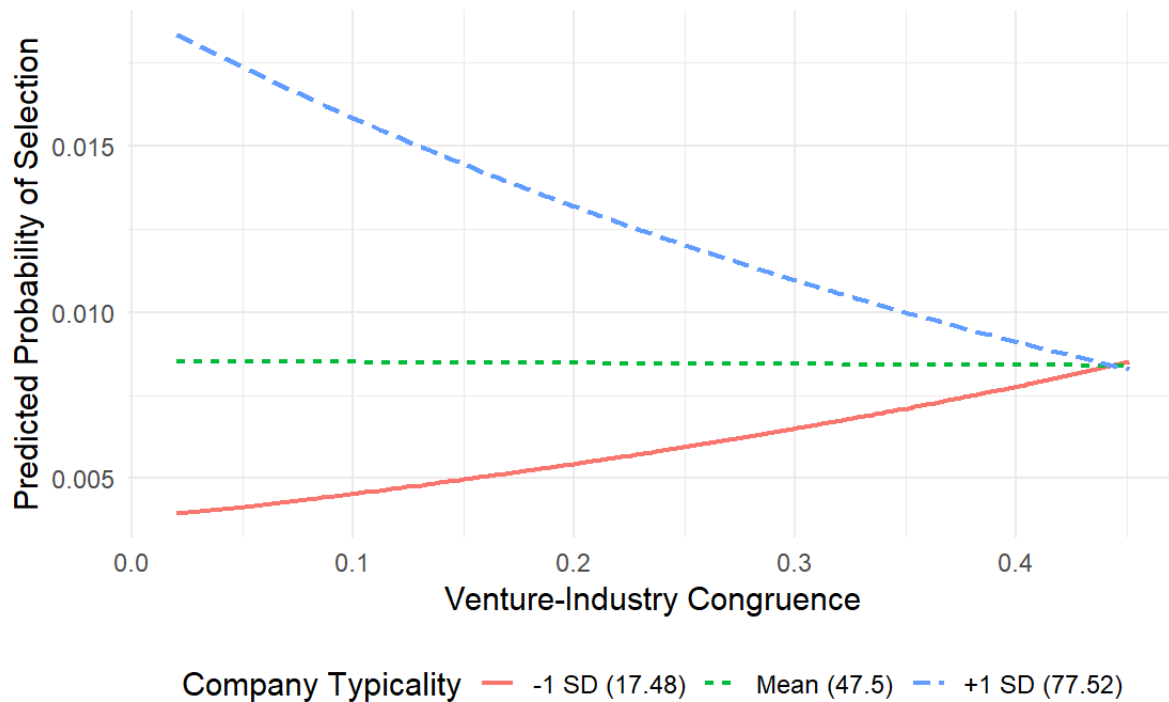


Figure 2: Predicted probabilities of selection for typical versus atypical firms across the range of venture–industry congruence in category-based searches.

The figure shows the pattern predicted by H2. When congruence is low, we observe the classic typicality premium: atypical firms are substantially less likely to be selected than typical firms. As congruence rises, this penalty for atypicality shrinks. At high congruence, the typicality effect is neutralized, and atypical firms reach parity with typical ones. The pattern illustrates that low-congruence contexts amplify the advantage of typical firms, whereas high-congruence contexts eliminate it. High congruence fosters leniency (Pontikes and Barnett 2015), lifting atypical firms to parity, but does not create a reversal that favors atypicality.

Robustness Checks

We conduct four additional checks to address concerns about search-result quality, company-level confounding, keyword-category overlap, and venture’s own typicality. Note that search field parameters (whether participants use specific search filters) affect the probability of company selection (Table 3). We interpret this as evidence that certain search parameters tend to return more relevant results, raising the concern that the observed effects of interest are not due to evaluation, but to the quality of the search results returned. To address this, we perform a stricter test on a specific subset of the data. Specifically, we analyze only those search result pages from which a participant selected exactly one competitor. This approach creates a sample of “equally successful” search outcomes, allowing us to compare

evaluations in contexts where the searcher demonstrably found what they were looking for. We re-estimate our main models (Table 3) on this subset ($n = 12,460$ for Model 4). The results, presented in Table A2 in the Online Supplement, are fully consistent with our main findings. The interaction between company typicality and venture-industry congruence remains significant and shows the same pattern of neutralization. This provides confidence that our results are not merely an artifact of varying search result quality but reflect a genuine shift in evaluation criteria.

A second concern might be that our results are driven by unobserved characteristics of the companies themselves, rather than their typicality within a given category. A more stringent test would be a model with company fixed effects, which would test whether the same company is evaluated differently depending on the congruence of the search context in which it appears. This approach is not feasible with our data, as it requires the same company to be evaluated by multiple participants under different conditions, a rare event in our behavioral simulation. Nonetheless, the main model addresses this concern in several ways. By controlling for direct venture-company semantic similarity, company size, and location, we account for the most salient observable characteristics that might otherwise confound the effect of typicality. This, combined with our use of participant-fixed effects, allows us to be confident that we are isolating the effect of a company's categorical positioning within the search context in which it appears.

Third, it is possible that participants used keywords that are the same as or similar to our industry labels. For instance, 7,073 of the total 76,814 companies in search results appeared when participants searched using keywords that directly matched an industry label. However, this means that some of the nonindustry search results are in fact industry-based, which would work against our H1. Our results are robust when re-estimating the main model without those 7,073 instances (Table A3 in the Online Supplement), suggesting that keyword-industry overlap does not drive our findings.

Fourth, high congruence may coincide with the searcher's own venture being typical of the focal category, so that entrepreneurs whose ventures already resemble the prototype stop using typicality to screen competitors. On this reading, neutralization would be driven by the venture's own typicality and not by goal-category fit. We re-estimated Model 4 adding an interaction between company typicality and venture typicality, alongside the congruence interaction. The added interaction is null, while the congruence interaction is unchanged (Table A4 in the Online Supplement). Neutralization tracks the fit between goal and category, not the venture's typicality.

Additional Analyses

Validating Category-Specific Engagement: Secondary Industry Typicality

A central claim of this paper is that typicality effects emerge from engaging a specific categorical structure, the focal industry structuring the search context. If they reflected a general preference for typical items, constant for a given audience

and context, then a company's positioning within any of its category memberships should affect selection. If our theory is correct, then only typicality relative to the focal category should matter, while typicality in other categories the company belongs to should not.

To test this mechanism directly, we examined whether typicality in a secondary (nonfocal) industry category also predicts company selection. For each search results page, we identified the second-most-frequent industry category represented among the displayed companies and calculated each company's typicality relative to this secondary industry using the same GPT-4o measurement approach described earlier. For participants who used the industry filter, the focal category is the industry they actively searched for, while the secondary category simply reflects additional memberships that some companies on the page happen to hold. We restricted the analysis to companies that are members of both the focal and secondary industries ($n = 39,021$ observations), as comparing typicality effects requires the company to actually belong to both categories being tested.

Table A5 in the Online Supplement presents three models testing this mechanism. Model 1 shows that when evaluated in isolation, secondary industry typicality predicts selection ($\beta = 0.003, p < 0.05$). However, Model 2 reveals a different pattern. When both focal and secondary typicalities are included, only focal industry typicality remains significant ($\beta = 0.005, p < 0.01$), while secondary typicality becomes nonsignificant. Model 3 provides direct evidence for the categorical engagement mechanism. Consistent with H1 and the main analysis, the interaction between focal industry typicality and category-based search is positive and significant ($\beta = 0.007, p < 0.05$), showing that focal typicality only matters when participants explicitly use industry categories to search. In contrast, the interaction between secondary industry typicality and category-based search is essentially zero, indicating no effect regardless of search method.

These results validate our core theoretical claim that typicality effects are not structural properties that operate uniformly across all categories a firm belongs to. Instead, they emerge through practices of categorical engagement and depend specifically on which category is activated as the organizing lens. For example, when participants search within the Education industry, a company's typicality *as an Education firm* matters for selection, but its typicality as a software firm (if software is also represented but less dominant) does not. The category that is explicitly engaged produces typicality effects. Categories that remain in the background, even when the firm belongs to them, do not shape evaluation.

Who Achieves Goal-Category Fit?

Having established that venture-industry congruence transforms typicality effects, we now explore whether the ability to achieve high goal-category congruence varies systematically across participants. Research on digital inequality shows that education and socioeconomic status predict internet skills and information-seeking effectiveness (DiMaggio et al. 2004; Hargittai and Hinnant 2008; Hargittai 2010). However, this work focuses primarily on general navigation proficiency and does not examine how actors engage formal categorical structures during search or the consequences for evaluation. We explore whether individual characteristics predict

Table 4: Antecedents of Venture-Industry Congruence.

	Venture-Industry Congruence
Education level	0.013 [†] (0.004)
Pre-existing idea	−0.031 [†] (0.006)
Gender male	−0.002 (0.006)
Age (10 years)	0.001 (0.004)
Constant	0.232 [†] (0.011)
Participant FE:	No
Observations	549
R^2	0.061
Adjusted R^2	0.054
Residual std. error	0.070 (df = 544)
F statistic	8.878 [†] (df = 4; 544)

Note: [†] $p < 0.01$.

effective categorical navigation, that is, identifying industry categories that align well with one's goals. This matters because achieving high congruence enables actors to identify atypical firms as competitors based on their merits rather than overlooking them due to categorical deviance.

We constructed a participant-level dependent variable: *Average venture-industry congruence*. For each search results page a participant generated, we calculated venture-industry congruence based on the most frequently appearing industry, as described previously. Each participant's average is the mean across all search pages they viewed. We estimate an ordinary least squares regression with participant-level predictors.

Educational level is a five-point ordinal scale ranging from 1 ("High school diploma") to 5 ("Doctoral degree"). Our sample was highly educated: 43 percent held bachelor's degrees and 76 percent held bachelor's degrees or higher. We treat this variable as continuous, excluding the 3 percent with missing data.

Pre-existing Venture concept is a binary variable coded as 1 if a participant reported that they had the venture concept before participating in the study (47 percent of sample). Participants coded as 0 reported coming up with their idea only when asked at the beginning of the survey. This variable may capture cognitive entrenchment. We include controls for *Gender* and *Age* (in years).

Table 4 presents the results. Educational level positively predicts average venture-industry congruence ($\beta = 0.013$, $p < 0.01$). Participants with more education generated results from more semantically relevant industries on average, a pattern in line with research on education and information-seeking quality.

Having a pre-existing idea negatively predicts average congruence ($\beta = -0.031$, $p < 0.01$). Participants who arrived with a pre-existing idea generated results from less semantically relevant industries on average. This pattern is consistent with

cognitive entrenchment limiting flexible exploration (Tripsas and Gavetti 2003). The demographic controls were not statistically significant.

These exploratory findings point to individual differences in goal-category congruence. These differences need not lie in searcher capability alone. Higher congruence among the more educated may reflect the kinds of ideas they pursue, and lower congruence among those with pre-existing ideas may reflect ideas held with greater conviction, or ideas that are more novel and category-spanning and so harder to map onto established industries. Whether such differences stem from the searcher or from the venture is a question for future research. Because achieving high congruence transforms which competitors an actor identifies, these differences, whatever their source, bear on who can identify innovative opportunities that defy simple category heuristics.

Discussion

Building on evidence of context-dependent typicality effects (Zuckerman et al. 2003; Ruef and Patterson 2009; Pontikes 2012; Goldberg et al. 2016), this paper identifies one mechanism through which such contingency operates: the search process preceding selection. Companies can be discovered through different means, and this search path determines the lens through which discovered items are assessed. We find that typicality within a category affects selection only when actors use that specific category to organize their search; typicality in other categories the firm belongs to becomes irrelevant. Furthermore, typicality effects are strongest when actors search in categories poorly aligned with their goals. In these low-congruence contexts, selection aims at identifying firms that best exemplify the category. When actors search in categories well-aligned with their goals, evaluation shifts to an item-level perspective where each firm is assessed on its individual merits regardless of its typicality within the category. Finally, tentative findings suggest that individuals differ in the goal-category congruence they achieve, shaping who becomes constrained to selecting typical category members and who can identify the most relevant options for their goals regardless of categorical structures.

Rethinking the Categorical Imperative

Our framework suggests that variation in typicality effects across contexts may arise not only from differences in audience preferences or institutional conditions, but also from differences in categorical engagement practices. Studies finding strong atypicality penalties (Zuckerman 1999; Hsu 2006; Kovács and Johnson 2014) may have examined settings where actors predominantly used category-based search and faced low goal-category congruence. Research documenting weaker effects or advantages to boundary spanning (Pontikes 2012; Leung 2014) may have observed contexts where actors either bypassed categorical structures during search or achieved high goal-category congruence. This account is complementary to multilogic views of valuation (Durand and Paoletta 2013; Glaser et al. 2020; Boulongne and Durand 2021; Gouvard and Durand 2023).

Our findings open several directions for research on market categories. First, they call for systematic investigation of variation in category knowledge. Prior

research has largely assumed that categorical positioning is perceived similarly by all audiences, with differences arising from different theories of value (Zuckerman 2017). However, having prior knowledge about a category is necessary to make category-based judgments (Hannan et al. 2019), and variations in that knowledge across audiences and individuals have received little attention. Different audiences may not know enough about institutionalized categorization systems to perceive atypicality, or may view items from perspectives where categorical dimensions lack sufficient weight. An amateur investor browsing industry lists may find typicality heuristics relevant, while expert investors using alternative search methods may not engage those categorical structures at all. Moreover, the same item can appear differently typical depending on evaluative context: a book that seems typical in a library may seem atypical on Amazon. Research should examine when and how categorical knowledge becomes activated rather than assuming it operates uniformly.

Second, future work should investigate properties of categorization systems themselves in terms of goal-category alignment. As Bowker and Star (2000) demonstrated, not all categorization systems are created equal. A medical taxonomy classifying snakes as “venomous or not” serves temperate-climate doctors but fails tropical doctors needing to know which antidote to use. Official industry classifications such as the North American Industry Classification System (NAICS) and Standard Industrial Classification (SIC) were developed for official statistical purposes, not for the diverse contemporary goals of entrepreneurs, investors, consumers, and employees. When do misalignments between taxonomic structure and typical user goals systematically bias evaluation toward prototypes? How do actors develop workarounds when established categories do not serve their purposes? Understanding categories as infrastructure means recognizing their effects depend on how well taxonomic partitions serve the purposes for which they are deployed.

Methodological Contributions

Our study demonstrates the value of behavioral simulation for capturing microprocesses that typically remain invisible to researchers. Where prior approaches infer categorical effects from outcomes, behavioral simulation lets researchers observe the process by which actors arrive at evaluations and identify the factors within that process that shape evaluation and selection. Categorical engagement and goal-category congruence are two such factors, but the method generalizes. It offers a way to study how features of the search and evaluation process come to determine outcomes that might otherwise be attributed to the properties of the items alone. The design also preserves agency that controlled experiments remove, since participants choose how to act, while capturing the specific goals that observational platform data lack. We see this as a productive approach for research wherever evaluation depends on a process that outcome data leave unobserved.

Our use of semantic similarity measures provides a productive tool for studying meaning-making in real time. By comparing venture descriptions with company descriptions using sentence transformers, we can quantify the concept of categorical fit, an aspect of classification that is difficult to capture with conventional measures of categorical membership or spanning. This approach makes it possible to examine the substantive alignment between actors’ goals and the categories they

use. The method could be applied to study other forms of cultural matching, from person-organization fit to brand-consumer alignment, opening new avenues for quantitative cultural analysis.

More broadly, our study provides a template for studying other forms of categorical engagement beyond search. The same methodology could examine how actors create new categories, modify existing ones, or strategically present themselves within categorical structures. By making the microprocesses of categorical engagement visible, behavioral simulation can illuminate how various social structures are activated, negotiated, and potentially transformed through practice.

Limitations and Boundary Conditions

Generalizability across Contexts

The primary boundary condition is whether our findings generalize beyond early competitor identification. The key premise of most organizational studies of categories is that they are driven by universal cognitive mechanisms (Hannan et al. 2019). This facilitates coherence of the field, as findings from restaurants (Kovács and Johnson 2014) inform studies in domains as different as venture investment (Pontikes 2012), and vice versa. Although competitor identification is not a standard empirical setting in research on categories, Cattani et al. (2017) view this setting as mutually informative with research on category-informed audience evaluation. Hence, we have consistently built our theory on insights from categorization research and expect that our findings generalize to audience evaluations. This being said, our participants approach the task as producers identifying their own competitors, and reliance on a categorical structure may differ for audiences such as critics or investors, especially those who have helped shape the categories they use. Examining how evaluative position conditions categorical engagement is a task for future research.

Several features of our empirical setting warrant consideration when assessing generalizability. First, identifying competitors for a new venture concept, especially one generated during the study, implies very high uncertainty about who the relevant competitors might be. Actors with deeper domain expertise may bypass categorical heuristics entirely. For instance, if we asked actual entrepreneurs to identify competitors on our platform, they might mostly search for competitors they already know, without reliance on categories. This suggests our findings may apply most directly to contexts where evaluators face similar uncertainty and information asymmetry, such as consumers exploring unfamiliar product categories, job seekers evaluating potential employers in new industries, or managers assessing acquisition targets in adjacent markets.

Second, the properties of the categorization system itself matter. Our study employed Crunchbase's industry taxonomy, which occupies a middle ground: not as institutionalized as official industry classifications like NAICS or SIC codes, but more established than user-generated folksonomies (Mathes 2004). In very well-established, mostly single-membership category systems like literary genres, we might not observe categorical engagement effects because there is no alternative to

searching with categories. Categorical framing would be constant. Conversely, in extremely fluid contexts with ad hoc or user-generated categories, everyone might navigate by custom paths and we would not observe systematic effects of engaging categories. Our findings apply most directly to intermediate cases where actors face choices about whether and how to engage established categorical structures.

A related question is whether categorical engagement operates similarly across different kinds of categorization systems. Our findings establish how congruence conditions typicality within industry categories, but the form this contingency takes may differ in other systems. Examining geographic location as the categorical structure, Kovács and Tyulyupo (2026) find a different form. There, congruence reverses the typicality effect, so that locally atypical firms become advantaged in high-congruence locations rather than merely losing their disadvantage. In our industry setting, high goal-category congruence neutralizes the typicality premium instead, with no sign that atypical firms gain a corresponding advantage. Congruence conditions typicality in both systems, but whether it neutralizes or reverses the effect may depend on the categorical dimension involved, leaving the generality of the neutralization pattern an open question.

Third, the search and selection platform shapes what can be observed. We employed an online information platform with multiple search pathways (industry filters, keywords, location, and size parameters). This design made categorical engagement variable and observable. In physical search contexts, such as libraries where books are organized by topics and genres with limited alternative access points, categorical engagement would be effectively mandatory. Similarly, when evaluators must rate items within specific categories, such as award judging panels or regulatory classification tasks, categorical engagement is imposed. Our findings about when actors engage categories and the consequences of that engagement apply to contexts where engagement is voluntary.

Observability and Measurement Challenges

In most real-world settings, observing the effects of categorical engagement and goal-category congruence presents substantial challenges. Consider cultural consumption: an individual's consideration set may be composite, assembled from friends' recommendations, algorithmic suggestions, and venues noticed while walking down the street. Search is protracted in time and multisource—consumers may use different search engines and information platforms, interleaving these with searches for other goals. Although our design enables clean observation of categorical engagement practices, it sacrifices the complexity and extended timescales of natural search processes.

Our measure of goal-category congruence remains subject to the problem of infinite dimensionality in category definitions, first raised by cognitive psychologists (Murphy and Medin 1985; Goldstone 1994) and viewed as relevant by organizational scholars (Durand and Paoletta 2013; Cattani et al. 2017). The challenge is that any two entities can be arbitrarily similar or dissimilar depending on which attributes are deemed relevant for comparison. Even the most advanced semantic similarity calculations do not fully resolve this challenge because semantically

distant companies might represent similar competitive threats from the perspective of a particular new venture. A mobile gaming venture might compete with a streaming service for consumers' leisure time despite their semantic distance in describing entertainment offerings. Our study does not overcome this deep-seated problem, but our measure of venture-industry congruence is grounded in direct venture-company similarity, which is the strongest predictor of competitor identification in our data. This suggests that semantic similarity, although imperfect, does capture meaningful dimensions of competitive relevance for the nascent entrepreneurs in our study.

Voluntary Engagement versus Imposed Categories

Our behavioral simulation preserves voluntary categorical engagement, allowing participants to choose whether and which categories to deploy. An alternative experimental approach would manipulate categorical framing by presenting identical firms with or without category labels, and by varying category-goal alignment. However, such a design addresses a different theoretical question. Our theory posits that typicality effects emerge when actors voluntarily choose to organize their search around categories. When categories are externally imposed, actors may not internalize categorical structures as organizing lenses, potentially weakening or eliminating the effects we observe. Whether imposed categorical framing produces similar patterns, or whether voluntary engagement is essential to our mechanisms, remains an important question for future research.

Conclusion

This study reveals that categorical structures shape markets through practices of engagement. The penalties for deviating from category norms emerge only when actors engage formal categories that misalign with their goals; when goals and categories align well, or when actors bypass categorical structures entirely, typicality becomes irrelevant. This engagement-based mechanism offers a process-level explanation for variation documented in prior work: why VCs using high-congruence search reward categorical ambiguity while consumers in low-congruence contexts penalize it (Pontikes 2012), why high-status actors escape penalties that constrain newcomers (Zuckerman et al. 2003), and why nascent systems fail to generate the penalties observed in mature ones (Ruef and Patterson 2009). Categories function as tools whose effects depend on whether and how they are deployed. Understanding when classification systems constrain evaluation requires attending to the microprocesses through which actors engage with them during decision-making.

Notes

- ¹ In Kovács and Tyulyupo (2026), participants act as consultants identifying competitors for a fixed venture concept, whereas here they identify competitors for their own venture ideas.

References

- Ahn, Jungsoo, Jean-Philippe Vergne, and Amanda Sharkey. 2026. "When Does Category Spanning Hurt or Help Producers?" *Strategic Management Journal* 47(7):1907–34. <https://doi.org/10.1002/smj.70079>
- Arora-Jonsson, Stefan, Nils Brunsson, Raimund Hasse, and Katarina Lagerström, eds. 2021. *Competition: What It Is and Why It Happens*. Oxford: Oxford University Press. <https://doi.org/10.1093/oso/9780192898012.001.0001>
- Barsalou, Lawrence W. 1983. "Ad Hoc Categories." *Memory & Cognition* 11(3):211–27. <https://doi.org/10.3758/BF03196968>
- Betancourt, Nathan, Inga J. Hoefer, and Filippo Carlo Wezel. 2025. "Atypicality and Accountability: Evidence from Five Experiments." *Organization Science* 36(2):838–61. <https://doi.org/10.1287/orsc.2022.16937>
- Botelho, Tristan L., Ranjay Gulati, and Olav Sorenson. 2024. "The Sociology of Entrepreneurship Revisited." *Annual Review of Sociology* 50:341–64. <https://doi.org/10.1146/annurev-soc-030222-014049>
- Boulongne, Romain and Rodolphe Durand. 2021. "Evaluating Ambiguous Offerings." *Organization Science* 32(2):257–72. <https://doi.org/10.1287/orsc.2020.1402>
- Bowker, Geoffrey C. and Susan Leigh Star. 2000. *Sorting Things Out: Classification and Its Consequences*. Cambridge, MA: MIT Press. <https://doi.org/10.7551/mitpress/6352.001.0001>
- Cattani, Gino, Joseph F. Porac, and Howard Thomas. 2017. "Categories and Competition." *Strategic Management Journal* 38(1):64–92. <https://doi.org/10.1002/smj.2591>
- Cattani, Gino, Daniel Sands, Joseph F. Porac, and Jason Greenberg. 2018. "Competitive Sensemaking in Value Creation and Capture." *Strategy Science* 3(4):632–57. <https://doi.org/10.1287/stsc.2018.0069>
- Chen, Ming-Jer. 1996. "Competitor Analysis and Interfirm Rivalry: Toward a Theoretical Integration." *Academy of Management Review* 21(1):100–34. <https://doi.org/10.2307/258631>
- Clark, Bruce H. 2011. "Managerial Identification of Competitors: Accuracy and Performance Consequences." *Journal of Strategic Marketing* 19(3):209–27. <https://doi.org/10.1080/0965254X.2011.557740>
- Clark, Bruce H. and David B. Montgomery. 1999. "Managerial Identification of Competitors." *Journal of Marketing* 63(3):67–83. <https://doi.org/10.1177/002224299906300305>
- Cutolo, Donato and Simone Ferriani. 2024. "Atypicality: Toward an Integrative Framework in Organizational and Market Settings." *Academy of Management Annals* 18(1):157–209. <https://doi.org/10.5465/annals.2022.0005>
- DiMaggio, Paul, Eszter Hargittai, Coral Celeste, and Steven Shafer. 2004. "Digital Inequality: From Unequal Access to Differentiated Use." Pp. 355–400 in *Social Inequality*, edited by Kathryn M. Neckerman. New York: Russell Sage Foundation.
- Douglas, Mary. 1986. *How Institutions Think*. Syracuse, NY: Syracuse University Press.
- Durand, Rodolphe and Lionel Paoletta. 2013. "Category Stretching: Reorienting Research on Categories in Strategy, Entrepreneurship, and Organization Theory." *Journal of Management Studies* 50(6):1100–23. <https://doi.org/10.1111/j.1467-6486.2011.01039.x>
- Dutton, Jane E. and Stephen A. Stumpf. 1991. "Using Behavioral Simulations to Study Strategic Processes." *Simulation & Gaming* 22(2):149–73. <https://doi.org/10.1177/1046878191222001>

- Ghose, Anindya, Panagiotis G. Ipeirotis, and Beibei Li. 2014. "Examining the Impact of Ranking on Consumer Behavior and Search Engine Revenue." *Management Science* 60(7):1632–54. <https://doi.org/10.1287/mnsc.2013.1828>
- Glaser, Vern L., Mariam Krikorian Atkinson, and Peer C. Fiss. 2020. "Goal-Based Categorization: Dynamic Classification in the Display Advertising Industry." *Organization Studies* 41(7):921–43. <https://doi.org/10.1177/0170840619883368>
- Goldberg, Amir, Michael T. Hannan, and Balázs Kovács. 2016. "What Does It Mean to Span Cultural Boundaries? Variety and Atypicality in Cultural Consumption." *American Sociological Review* 81(2):215–41. <https://doi.org/10.1177/0003122416632787>
- Goldstone, Robert L. 1994. "The Role of Similarity in Categorization: Providing a Groundwork." *Cognition* 52(2):125–57. [https://doi.org/10.1016/0010-0277\(94\)90065-5](https://doi.org/10.1016/0010-0277(94)90065-5)
- Gouvard, Paul and Rodolphe Durand. 2023. "To Be or Not to Be (Typical): Evaluation-Mode Heterogeneity and Its Consequences for Organizations." *Academy of Management Review* 48(4):659–80. <https://doi.org/10.5465/amr.2020.0314>
- Gur, Furkan Amil and Thomas Greckhamer. 2019. "Know Thy Enemy: A Review and Agenda for Research on Competitor Identification." *Journal of Management* 45(5):2072–100. <https://doi.org/10.1177/0149206317744250>
- Hannan, Michael T., Gaël Le Mens, Greta Hsu, Balázs Kovács, Giacomo Negro, László Pólos, Elizabeth Pontikes, and Amanda J. Sharkey. 2019. *Concepts and Categories: Foundations for Sociological and Cultural Analysis*. New York: Columbia University Press. <https://doi.org/10.7312/hann19272>
- Hargittai, Eszter. 2010. "Digital Na(t)ives? Variation in Internet Skills and Uses among Members of the 'Net Generation'." *Sociological Inquiry* 80(1):92–113. <https://doi.org/10.1111/j.1475-682X.2009.00317.x>
- Hargittai, Eszter and Amanda Hinnant. 2008. "Digital Inequality: Differences in Young Adults' Use of the Internet." *Communication Research* 35(5):602–21. <https://doi.org/10.1177/0093650208321782>
- Hsu, Greta. 2006. "Jacks of All Trades and Masters of None: Audiences' Reactions to Spanning Genres in Feature Film Production." *Administrative Science Quarterly* 51(3):420–50. <https://doi.org/10.2189/asqu.51.3.420>
- Hsu, Greta, Michael T. Hannan, and Özgecan Koçak. 2009. "Multiple Category Memberships in Markets: An Integrative Theory and Two Empirical Tests." *American Sociological Review* 74(1):150–69. <https://doi.org/10.1177/000312240907400108>
- Katila, Riitta, Eric L. Chen, and Henning Piezunka. 2012. "All the Right Moves: How Entrepreneurial Firms Compete Effectively." *Strategic Entrepreneurship Journal* 6(2):116–32. <https://doi.org/10.1002/sej.1130>
- Kim, Bo Kyung, and Michael Jensen. 2011. "How Product Order Affects Market Identity: Repertoire Ordering in the U.S. Opera Market." *Administrative Science Quarterly* 56(2):238–56. <https://doi.org/10.1177/0001839211427535>
- Kovács, Balázs, Gianluca Carnabuci, and Filippo Carlo Wezel. 2021. "Categories, Attention, and the Impact of Inventions." *Strategic Management Journal* 42(5):992–1023. <https://doi.org/10.1002/smj.3271>
- Kovács, Balázs and Michael T. Hannan. 2010. "The Consequences of Category Spanning Depend on Contrast." Pp. 175–201 in *Categories in Markets: Origins and Evolution*, edited by Greta Hsu, Giacomo Negro, and Özgecan Koçak. Research in the Sociology of Organizations, vol. 31. Bingley, UK: Emerald Group Publishing. [https://doi.org/10.1108/S0733-558X\(2010\)0000031008](https://doi.org/10.1108/S0733-558X(2010)0000031008)

- Kovács, Balázs and Michael T. Hannan. 2015. "Conceptual Spaces and the Consequences of Category Spanning." *Sociological Science* 2:252–86. <https://doi.org/10.15195/v2.a13>
- Kovács, Balázs and Rebeka Johnson. 2014. "Contrasting Alternative Explanations for the Consequences of Category Spanning: A Study of Restaurant Reviews and Menus in San Francisco." *Strategic Organization* 12(1):7–37. <https://doi.org/10.1177/1476127013502465>
- Kovács, Balázs and Alex Tyulyupo. 2026. "Cognitive Cartography: How Geographic Categories and Firm Location-Typicality Shape Competitor Identification." *Industrial and Corporate Change*. <https://doi.org/10.1093/icc/dtag031>
- Le Mens, Gaël, Balázs Kovács, Michael T. Hannan, and Guillem Pros. 2023a. "Uncovering the Semantics of Concepts Using GPT-4." *Proceedings of the National Academy of Sciences* 120(49):e2309350120. <https://doi.org/10.1073/pnas.2309350120>
- Le Mens, Gaël, Balázs Kovács, Michael T. Hannan, and Guillem Pros. 2023b. "Using Machine Learning to Uncover the Semantics of Concepts: How Well Do Typicality Measures Extracted from a BERT Text Classifier Match Human Judgments of Genre Typicality?" *Sociological Science* 10:82–117. <https://doi.org/10.15195/v10.a3>
- Leung, Ming D. 2014. "Dilettante or Renaissance Person? How the Order of Job Experiences Affects Hiring in an External Labor Market." *American Sociological Review* 79(1):136–58. <https://doi.org/10.1177/0003122413518638>
- Mathes, Adam. 2004. "Folksonomies: Cooperative Classification and Communication through Shared Metadata." Unpublished manuscript, Graduate School of Library and Information Science, University of Illinois at Urbana-Champaign.
- McCall, Morgan W. and Michael M. Lombardo. 1982. "Using Simulation for Leadership and Management Research: Through the Looking Glass." *Management Science* 28(5):533–49. <https://doi.org/10.1287/mnsc.28.5.533>
- Murphy, Gregory L. and Douglas L. Medin. 1985. "The Role of Theories in Conceptual Coherence." *Psychological Review* 92(3):289–316. <https://doi.org/10.1037/0033-295X.92.3.289>
- Pontikes, Elizabeth G. 2012. "Two Sides of the Same Coin: How Ambiguous Classification Affects Multiple Audiences' Evaluations." *Administrative Science Quarterly* 57(1):81–118. <https://doi.org/10.1177/0001839212446689>
- Pontikes, Elizabeth G. and William P. Barnett. 2015. "The Persistence of Lenient Market Categories." *Organization Science* 26(5):1415–31. <https://doi.org/10.1287/orsc.2015.0973>
- Porac, Joseph F. and Howard Thomas. 1990. "Taxonomic Mental Models in Competitor Definition." *Academy of Management Review* 15(2):224–40. <https://doi.org/10.2307/258155>
- Porac, Joseph F., Howard Thomas, and Charles Baden-Fuller. 2011. "Competitive Groups as Cognitive Communities: The Case of Scottish Knitwear Manufacturers Revisited." *Journal of Management Studies* 48(3):646–64. <https://doi.org/10.1111/j.1467-6486.2010.00988.x>
- Porter, Michael E. 1980. "Industry Structure and Competitive Strategy: Keys to Profitability." *Financial Analysts Journal* 36(4):30–41. <https://doi.org/10.2469/faj.v36.n4.30>
- Reimers, Nils and Iryna Gurevych. 2019. "Sentence-BERT: Sentence Embeddings Using Siamese BERT-Networks." *arXiv:1908.10084*. <https://arxiv.org/abs/1908.10084>
- Rosch, Eleanor and Carolyn B. Mervis. 1975. "Family Resemblances: Studies in the Internal Structure of Categories." *Cognitive Psychology* 7(4):573–605. [https://doi.org/10.1016/0010-0285\(75\)90024-9](https://doi.org/10.1016/0010-0285(75)90024-9)

- Ruef, Martin and Kelly Patterson. 2009. "Credit and Classification: The Impact of Industry Boundaries in Nineteenth-Century America." *Administrative Science Quarterly* 54(3):486–520. <https://doi.org/10.2189/asqu.2009.54.3.486>
- Sands, Daniel, Gino Cattani, Joseph F. Porac, and Jason Greenberg. 2021. "Competition as Sense Making." Pp. 26–47 in *Competition: What It Is and Why It Happens*, edited by Stefan Arora-Jonsson, Nils Brunsson, Raimund Hasse, and Katarina Lagerström. Oxford: Oxford University Press. <https://doi.org/10.1093/oso/9780192898012.003.0002>
- Sgourev, Stoyan V. and Niek Althuisen. 2014. "'Notable' or 'Not Able': When Are Acts of Inconsistency Rewarded?" *American Sociological Review* 79(2):282–302. <https://doi.org/10.1177/0003122414524575>
- Simon, Herbert A. 1955. "A Behavioral Model of Rational Choice." *Quarterly Journal of Economics* 69(1):99–118. <https://doi.org/10.2307/1884852>
- Tripsas, Mary and Giovanni Gavetti. 2003. "Capabilities, Cognition, and Inertia: Evidence from Digital Imaging." Pp. 393–412 in *The SMS Blackwell Handbook of Organizational Capabilities: Emergence, Development, and Change*, edited by Constance E. Helfat. Malden, MA: Blackwell Publishing. <https://doi.org/10.1002/9781405164054.ch23>
- Tversky, Amos. 1977. "Features of Similarity." *Psychological Review* 84(4):327–52. <https://doi.org/10.1037/0033-295X.84.4.327>
- Vogel, Tobias, Moritz Ingendahl, and Piotr Winkielman. 2021. "The Architecture of Prototype Preferences: Typicality, Fluency, and Valence." *Journal of Experimental Psychology: General* 150(1):187–94. <https://doi.org/10.1037/xge0000798>
- Winkielman, Piotr, Jamin Halberstadt, Tedra Fazendeiro, and Steve Catty. 2006. "Prototypes Are Attractive Because They Are Easy on the Mind." *Psychological Science* 17(9):799–806. <https://doi.org/10.1111/j.1467-9280.2006.01785.x>
- Zuckerman, Ezra W. 1999. "The Categorical Imperative: Securities Analysts and the Illegitimacy Discount." *American Journal of Sociology* 104(5):1398–438. <https://doi.org/10.1086/210178>
- Zuckerman, Ezra W. 2017. "The Categorical Imperative Revisited: Implications of Categorization as a Theoretical Tool." Pp. 31–68 in *From Categories to Categorization: Studies in Sociology, Organizations and Strategy at the Crossroads*, edited by Rodolphe Durand, Nina Granqvist, and Anna Tyllström. Research in the Sociology of Organizations 51. Bingley, UK: Emerald Publishing. <https://doi.org/10.1108/S0733-558X20170000051001>
- Zuckerman, Ezra W., Tai-Young Kim, Kalinda Ukanwa, and James von Rittmann. 2003. "Robust Identities or Nonentities? Typecasting in the Feature-Film Labor Market." *American Journal of Sociology* 108(5):1018–73. <https://doi.org/10.1086/377518>

Acknowledgments: We thank Ceclin Begbie and Robert Bartholomew for their management of the behavioral lab. We thank Maciej Workiewicz, Elisa Operti, and Stoyan Sgourev for their input. We gratefully acknowledge financial support from the Yale School of Management. We benefited from feedback received at the 2025 Nagymaros Conference and the Yale School of Management OBID seminar.

Alex Tyulyupo: Yale School of Management, Yale University.
E-mail: alex.tyulyupo@yale.edu.

Balázs Kovács: Yale School of Management, Yale University.
E-mail: balazs.kovacs@yale.edu.