

A deep-learning model of prescient ideas demonstrates that they emerge from the periphery

Paul Vicinanza ^{a,*}, Amir Goldberg ^a and Sameer B. Srivastava ^b

^aGraduate School of Business, Stanford University, 655 Knight Way, Stanford, CA 94305, USA

^bHaas School of Business, University of California, Berkeley, 2220 Piedmont Ave, Berkeley, CA 94720, USA

*To whom correspondence should be addressed: Email: pvicinan@stanford.edu

Edited By: Jay Van Bavel

Abstract

Where do prescient ideas—those that initially challenge conventional assumptions but later achieve widespread acceptance—come from? Although their outcomes in the form of technical innovation are readily observed, the underlying ideas that eventually change the world are often obscured. Here, we develop a novel method that uses deep learning to unearth the markers of prescient ideas from the language used by individuals and groups. Our language-based measure identifies prescient actors and documents that prevailing methods would fail to detect. Applying our model to corpora spanning the disparate worlds of politics, law, and business, we demonstrate that it reliably detects prescient ideas in each domain. Moreover, counter to many prevailing intuitions, prescient ideas emanate from each domain's periphery rather than its core. These findings suggest that the propensity to generate far-sighted ideas may be as much a property of contexts as of individuals.

Significance Statement:

Prescient ideas—those that initially challenge conventional assumptions but later achieve widespread acceptance—drive the most consequential changes in society. Yet such ideas are difficult to directly observe, leading researchers to focus on how they manifest in the form of domain-specific technical innovations. Here, we develop a novel, deep-learning method to identify prescient ideas in everyday language. We demonstrate that this common approach to quantifying prescience can be applied across such disparate domains as law, politics, and business and that prescient ideas arise from the periphery rather than the center of these fields. Although the capacity to generate prescient ideas is often attributed to individuals, its origins may have as much to do with place as with people.

Introduction

Where do prescient ideas—those that initially challenge conventional assumptions but later achieve widespread acceptance—come from? Prior research, predominately focused on the scientific and technical realms, traces these ideas to brilliant individuals (1, 2) or the characteristics of inventors and their teams (3–7). Because ideas themselves are difficult to observe, this literature focuses on how they manifest in the form of tangible artifacts such as patents or scientific publications. Yet most ideas, from disruptive business strategies to paradigm-shifting legal interpretations, do not translate into discrete artifacts that are amenable to such analysis. Instead, they fundamentally transform the prevailing assumptions and beliefs in a domain—but often in subtle ways.

In the realm of politics, for instance, legislators regularly introduce and contest novel ideas that later become taken-for-granted assumptions. In the debates that raged about civil rights legislation in the United States during the 1960s, two of the staunchest opponents of these bills were Senators John Stennis and James Eastland. Although both voted against every major piece of civil

rights legislation, they diverged in the nature of the opposition they put forward. Whereas Eastland—the least prescient senator according to our model—framed his opposition using overtly racist arguments (see Supplementary Material and Figure 1), Stennis, the senator deemed most prescient by our model, was among the first to base his objections on the principles of “color blindness,” limited government, and individual freedom (8). This more muted and indirect set of arguments would later become commonplace among opponents of civil rights legislation, laying the groundwork for contemporary conservative talking points on race relations in the United States. Stennis’ strategy for reframing race relations in the United States cannot be simply reduced to a single statement or artifact. Nor was it merely about novel rhetoric. Instead, it involved discourse that sought to deftly rearrange the tacit assumptions about the interrelationships between race, policy, and conservative values.

Building on this intuition, we propose that prescient ideas are incommensurable, in the Kuhnian sense (9), with conventional logic. Kuhn, in his groundbreaking work *The Structure of Scientific Revolution*, argued that measuring paradigm shifts requires

Competing Interest: The authors declare no competing interest.

Received: July 18, 2022. **Accepted:** December 2, 2022

© The Author(s) 2022. Published by Oxford University Press on behalf of the National Academy of Sciences. This is an Open Access article distributed under the terms of the Creative Commons Attribution-NonCommercial-NoDerivs licence (<https://creativecommons.org/licenses/by-nc-nd/4.0/>), which permits non-commercial reproduction and distribution of the work, in any medium, provided the original work is not altered or transformed in any way, and that the work is properly cited. For commercial re-use, please contact journals.permissions@oup.com

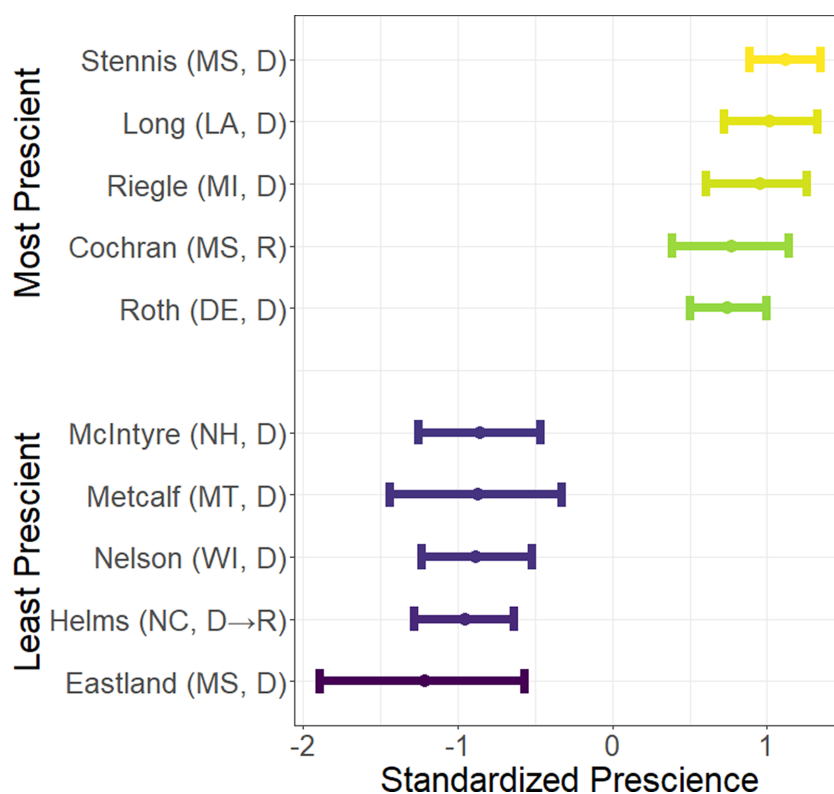


Fig. 1. Prescient senators, minimum three congressional terms. Mean prescience is computed, standardized, and bootstrapped 10k times at the politician-quarter level.

“indirect and behavioral evidence” and suggested that a promising source of such evidence is the language scientists use to describe the world. He gives the example of the word “planet” during the Copernican Revolution. By stripping the Sun of the designation of “planet,” Copernicans reconceptualized all celestial bodies and the relationships between them. Thus, shifts in language can sometimes signal deeper changes in cognitive understanding among actors in a domain, leaving linguistic markers—such as the use of “planet” in a somewhat modified sense—that foretell fundamental changes in the world.

Groundbreaking advances in natural language processing and deep learning have made it possible to extract such markers from the natural language that people use. Word embedding models capture culturally salient dimensions by representing words in a high-dimensional space. Contextual embedding models, which explicitly account for the context of each word, allow these dimensions to vary across contexts (10–13). We identify the linguistic markers of prescient ideas using Bidirectional Encoder Representations from Transformers (BERT) (14), a deep neural network that encodes the semantic and contextual information of language. Taking advantage of BERT’s ability to predict words given their context, we first develop a measure of *contextual novelty*: utterances that are poorly predicted by the model. To measure *prescience*, we compare the contextual novelty at the time ideas are first expressed within a domain to contextual novelty at a later point in time. Ideas that are incommensurable with conventional logic at the time of their expression but become more commensurable in the future are deemed prescient by our model. Importantly, this approach is agnostic to whether prescient individuals or groups are themselves the inventor of an idea, are early adopters, or are intentional prognosticators.

Our novel method enables us to unearth prescient ideas from natural language independent of the form in which the idea is ultimately realized and across a broad range of domains. Moreover, it allows us to answer a fundamental question in the science of ideas and innovation: Where do prescient ideas come from? Existing theories and evidence lead to competing expectations about whether prescient ideas will emerge from the periphery or the core of a given domain. On one hand, prior work on technological innovation suggests that transformative ideas emerge from the periphery because actors on the outskirts are less bound by institutional constraints and have greater freedom to explore new ideas (15–17). Yet in many nontechnical domains such as law, far-sighted ideas are commonly assumed to emerge from sage and established actors such as Supreme Court justices (18). Such actors are believed to have the resources and experience needed to rethink prevailing assumptions and the power to implement their vision. Even in the realm of technological innovation, prescient ideas often seem to emanate from the core, rather than peripheral, positions. Apple’s groundbreaking iPhone, for example, was released when it was already among the world’s most dominant technology companies. We attribute these mixed findings and contradictory intuitions to the fact that prior work has focused on studying the tangible artifacts that result from prescient ideas rather than interrogating where the ideas themselves came from.

To address this gap, we apply our method for detecting the linguistic markers of prescient ideas to three corpora that span the disparate domains of politics, law, and business. We first validate our measure by establishing that prescience is recognized and rewarded in a given field: Prescient court decisions are more highly cited, prescient politicians ascend to more powerful committees, and prescient firms have higher future market valuations.

Delving into the origins of prescient ideas, we next demonstrate that highly prescient ideas emanate from the periphery rather than the center of the field.

Defining prescient ideas

What makes an idea prescient? One key ingredient would appear to be novelty, as prescient ideas depart from the accepted conventions of the moment. Yet abundant prior research has shown that novelty does not by itself guarantee success (19–21). We propose that prescient ideas have two essential properties. First, they are novel in a particular way: They rethink the contextual assumptions that predominate a given field. By contextual assumptions, we mean those that (a) are central to a domain's logics of action and (b) guide a set of interdependent choices about how to configure activities for success in the field. In 1970, for example, Congress passed the Racketeer Influenced and Corrupt Organizations Act (RICO) to target organized criminal enterprises—in particular the Mafia. A small group of imaginative prosecutors soon seized upon RICO's ambiguous language to prosecute such wide-ranging civil crimes as mail fraud and stock manipulation (22). This approach was contextually novel in that it applied existing statutes intended for one set of actors to an entirely different class of actors and criminal activities.

Second, prescient ideas foreshadow how the domain will evolve in the future. While RICO's scope was initially limited, this approach to interpreting and applying RICO statutes well beyond their original scope became commonplace among prosecutors and judges. It has since been used to prosecute organizations ranging from the Catholic Church to Major League Baseball to British Petroleum (BP) following the Deepwater Horizon oil spill (22). Notice that these two ingredients—contextual novelty when an idea is initially articulated and widespread acceptance in its future—can appear in an individual's discourse even when the person does not explicitly set out to predict the future, influence others, or even change the world. Moreover, prescient ideas can only be detected after the fact—that is, once the future state of the world is known.

Developing a language-based measure of prescience

Ideas are, of course, hatched by individuals and often expressed in discourse. The core idea of our approach is to measure the extent to which ideas expressed in routine discourse are linguistic markers of prescient ideas. Our technique relies on the intuition that prescient ideas depart from prevailing ideas at the time of introduction but become commonplace in the future (5, 23). In particular, we use BERT, which learns the semantic and syntactic structure of language (in part) through a masked-word prediction task (14). BERT repeatedly predicts different masked (hidden) words in a sentence given the rest of the sentence, with the aim of minimizing the cross-entropy loss between the predicted and actual word. While most researchers apply BERT's model architecture to solve downstream tasks such as machine translation or text classification, we use the probabilistic features of the model to assess the extent to which a given set of ideas are prescient in their field.

Similar to how prior work trains separate word embedding models on a temporally split corpus to uncover semantic shifts in word meaning (10), by training separate BERT models over a split corpus, we reveal how the likelihood of specific words, phrases, and sentences evolves over time. Following standard practice, we

begin with a pretrained model and then fine-tune it to a given time interval (e.g., a year or presidential term) in each of our domains of interest.

To measure prescience, we begin by considering perplexity, the exponentiated cross-entropy loss, which can be intuitively understood as the inverse-likelihood of the model generating a word or a document (normalized by the number of words). Higher perplexity scores correspond to unusual or unexpected utterances. We define *contextual novelty*, $CN(s)$, as the product of word-level perplexities in a sentence, s , normalized by the number of words in the sentence. We use mean sentence-level perplexity values to derive a measure of contextual novelty at the document level. *Prescience* can then be operationalized as the percentage decrease in *contextual novelty* of a sentence between two time periods. Figure 2 provides a schematic representation of our measurement approach, using a sentence from our legal data that was deemed highly prescient by our model. Depending on the context, we can then aggregate our measure of prescient discourse at different levels of analysis. In some settings, we can identify individuals who are apt to express prescient ideas. These individual-level measures can be aggregated to the level of social groups or organizations that might be more salient in other contexts. In still other settings, the relevant unit of analysis might be a prescient document.

This approach to measuring prescience offers several advantages over prior work such as topic models, n -grams, or term frequency-inverse document frequency (TF-IDF) vectorization (5, 23, 24). When applied to scientific innovation, these methods are commonly used to identify the introduction of new terms or bigrams. This assumes that novel ideas necessarily translate into new terminology. Yet innovative ideas in nontechnical domains, such as those studied in this paper, are rarely accompanied by new terms. Instead, prescient ideas emerge as actors reconceptualize existing topics in incommensurate ways. Both Stennis and Eastland ardently denounced racial desegregation, yet only Stennis is considered prescient by framing the issue around color-blindness and limited government. By taking context into account, contextual embeddings capture such nuance.

Another feature of our approach is that any potential biases toward high-perplexity sentences—such as rare tokens or errors in optical character recognition (OCR)—are netted out in the numerator. Likewise, discussions unrelated to prescience are netted out because the likelihood of a sentence must shift over time to result in a nonzero contribution to prescience. Traditional preprocessing steps, such as removing stop words and punctuation, stemming, or converting to lowercase, are unnecessary in contextual embedding models. Unlike topic models, our approach does not require tuning hyperparameters—though the researcher does have to make choices about how to partition the data. For more details, please refer to the Supplementary Material.

Results

Empirical settings

We apply this method to identify prescient ideas in the domains of politics (4.9 million floor speeches given by members of the US House of Representatives and the US Senate), law (4.2 million rulings on US State and Federal cases), and business [108,334 quarterly earnings calls (QECs) in which the management teams of publicly traded firms lay out their vision and strategy for the

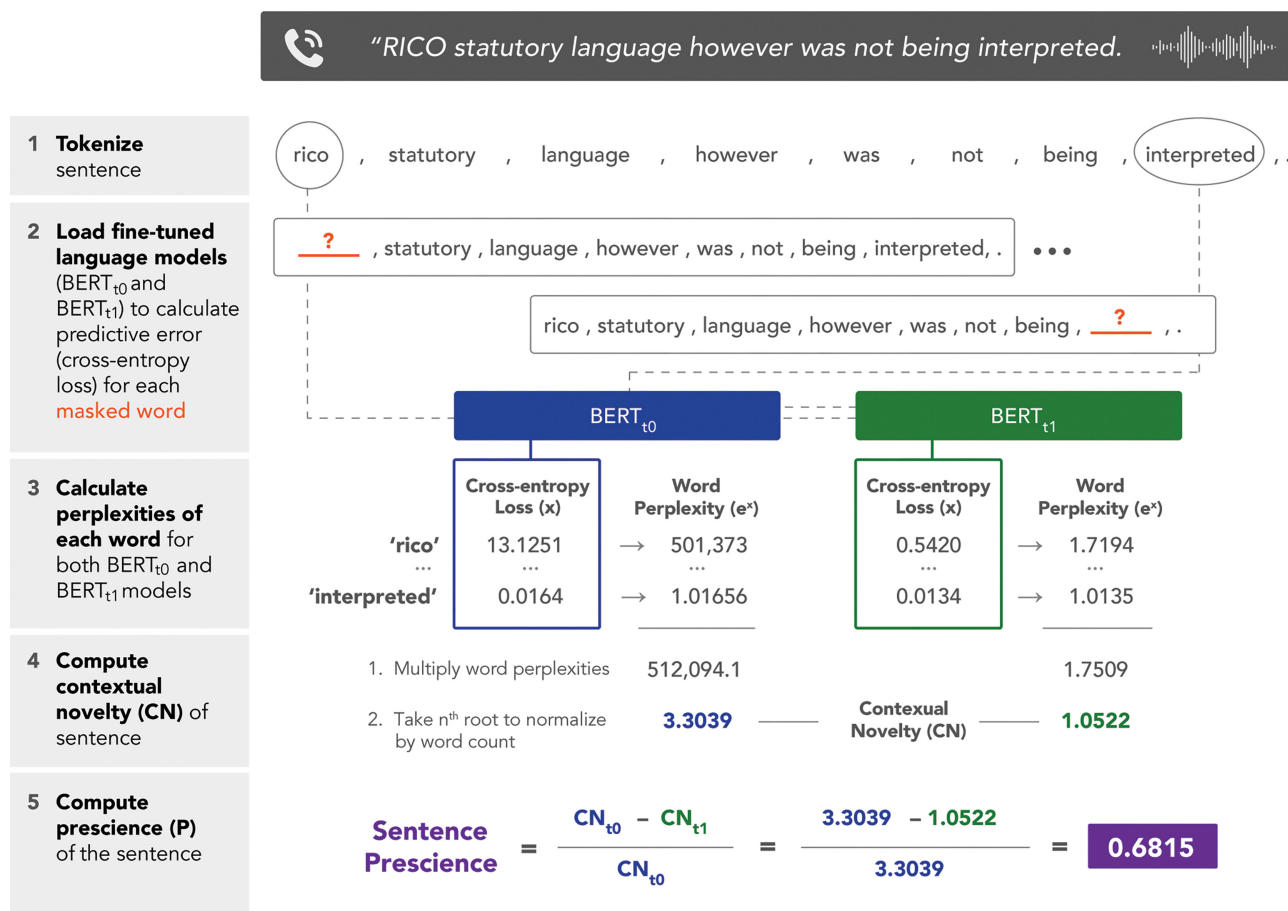


Fig. 2. Illustration of how prescience is computed based on a sentence from the legal dataset that the model deems highly prescient. This sentence rates highly in prescience because the RICO token better predicted the future period, when RICO's statutory language was heavily contested by the courts (9).

company to financial analysts who cover their stock]. Given that the corpora vary considerably in the time periods they cover and the nature of the discourse they include, we use slightly different approaches to fine-tuning BERT and defining the salient time horizon across the three (Supplementary Material Appendix).

Model validation

First, we return to the case studies of the two senators who were identified by our model as most and least prescient in the political dataset (Figure 1): Senators John Stennis and James Eastland, respectively. While both Stennis and Eastland were segregationists from Mississippi, Stennis is remembered as a conscientious leader of the Senate, while Eastland now symbolizes southern intransigence on racial integration. Scholars have traced these divergent judgments to Stennis' pioneering use of novel arguments—such as appeals to “color blindness,” limited government, and individual freedom—to oppose civil rights legislation (8). Our model echoes this conclusion, rating Stennis' 1964 utterance, “I will join hands with any Senator in trying to devise a voluntary, noncoercive, and noncompulsive plan under which both races can work and live together in harmony and make progress,” as being in the 99th percentile in prescience. Our method also correctly anticipates the impending disappearance of Eastland's overt brand of bigotry: his 1963 statement against civil rights legislation, “can he rationally write that races are equally law-abiding when the police statistics jeer at the thought?” is more than two standard de-

viations below mean prescience. Other senators in the most/least prescient list add further credence to our measure. Russell B. Long, the second most prescient Senator, chaired the Senate Finance Committee for over 25 y and is widely regarded as one of the most powerful, respected, and influential Senators of the 20th century (25). In contrast, the second least prescient senator, Jesse Helms, ardently opposed non-white, disability, gay, and feminist rights and is perhaps the most conservative politician of the modern era.

To illustrate the granularity and efficacy of our methodological approach, we examine the most and least prescient uses of the word “race” in Table 1. Contextual embeddings compute perplexities at the level of individual words, and we can examine instances when the word “race” is used in ways that the model deems highly prescient (i.e., the focal word is better predicted by a model trained on text from future periods relative to a model trained on the contemporaneous period) versus minimally prescient. Conducting this analysis for text from the 87th Congress (1961–1963), we observe stark differences between how the term is used in its most versus least prescient forms. Highly prescient uses are uttered in the context of discourse that seeks to eradicate racial discrimination and that denounces race-bating argumentation. The least prescient uses seem strikingly anachronistic by today's standards—for example, using terminology such as the “distinguished members of the white race” and “race suicide.”

To more systematically validate our measurement strategy, we conduct an additional analysis based on the so-called “Gingrich Senators,” a group of 33 Republican members of the US House of

Table 1. Most and least prescient uses of the word “race.”

Most prescient	I am astonished at the crude attempt of the president to inject a race issue into what is essentially a matter of public policy.
	This is an attempt to exploit Mr. Weavers name and race to gain a purely partisan end.
	The issue was a race issue pure and simple.
	Let us take the necessary legislative action to leave no room for doubt that we will no longer tolerate in the capital of the United States discrimination in real estate transactions based upon a race issue.
	The new Committee on Equal Employment Opportunity is moving to enforce executive orders against race discrimination in hiring workers on work done under Government contract.
Least prescient	So my friends, the white race is on trial very seriously on trial.
	Let me say to the gentleman in my familiarity with his record and the legislation he has introduced there has never been anything that was for the benefit of the white race.
	There is no race suicide in the great State of Michigan.
	Already the conflict in Algeria has degenerated into race war.
	I think I hardly need to add that many prominent and distinguished members of the white race are members of the organization on its board of directors and in key positions on its various committees.

The most and least prescient uses of the word “race” by politicians in the 87th Congress (1961–1963). Examples are restricted to the context of racial groups—that is, excluding sentences discussing such phrases as “space race” or “political race.”

Representatives who, according to scholarship in political science (26), served with Newt Gingrich in the House and espoused a new style of highly partisan, divisive, and obstructionist politics. These Representatives were later elected to the US Senate, to which they brought an overtly antagonistic style of discourse. The Gingrich Senators represent a textbook example of prescience: They were a group of individuals who used language that was indicative of a new brand of politics, which later became institutionalized. If our strategy for measuring prescience is valid, it should rate these 33 individuals as being more prescient than other Republicans of a comparable vintage. As Figure 3 indicates, this is indeed the case.

Next, we consider the relationship between prescience and success in a given domain. If our measure is indeed capturing successful ideas before their success, then first-moving actors who express these ideas should, on average, be recognized and rewarded by the relevant audiences in their field. Consistent with this expectation (Figure 4), prescience is positively related to: a politician’s likelihood of being reelected and her status attainment (Panels A and B; Table S1); a legal ruling’s total citations and probability of being a landmark ruling (Panels C and D; Table S2); and a firm’s annual stock returns (Panel E; Table S3). Indeed, it is only the highly prescient firms (top 5%) that achieve breakthrough levels of cumulative returns (Panel F).

The origins of prescient ideas

We turn next to investigating the sources of prescient ideas. Positions in a given domain can be thought of as varying along a continuum from the core, which tends to be occupied by established actors that shape the institutions and norms of a

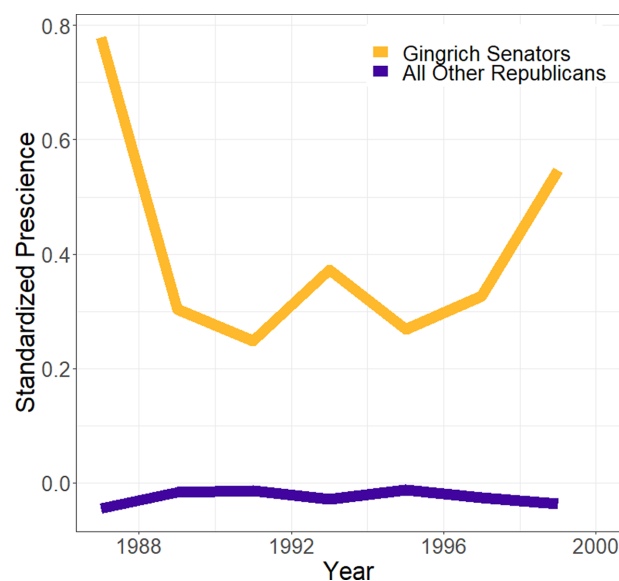


Fig. 3. Gingrich senators. Average standardized prescience for the Gingrich senators and all other Republicans by congressional term.

field, to the periphery, which is typically populated by upstart actors that have less ability to coordinate or influence behavior (15, 27). Relative to core actors, those in the periphery typically face fewer institutional constraints and are therefore more likely to generate novel ideas (17). Although the core/periphery distinction has been widely invoked, it has been operationalized in different ways across empirical contexts. Given that our three empirical settings are quite different from one another, we similarly develop measures of core versus peripheral positions that are specific to and appropriate for each context. These choices are discussed in greater detail in the Supplementary Material.

Figure 5 shows that, across all three settings, highly prescient ideas—those at or above the 95th percentile of our continuous measure of prescience—emanate from the periphery rather than the core. In politics, eigenvector network centrality (Panel A), k-core network centrality (Panel B), degree centrality, and closeness centrality are negatively related to the likelihood of a politician demonstrating elite prescience (Table S5). In law, lower (more peripheral) courts are more likely to produce highly prescient decisions than upper (more belonging to the core) courts (Panel C; Table S6). Indeed, the US Supreme Court, despite having the highest status and the power to set precedent, exhibits the lowest rate of highly prescient decisions. Highly prescient decisions are 22 times more likely to come from the State Appeals Courts than the US Supreme Court. Estimating models with judge fixed effects that take advantage of the fact that judges are sometimes promoted across the judicial status hierarchy, we find that the likelihood of a judge authoring a highly prescient ruling declines by 0.8% points after she is promoted from the US District Court—the lowest rung of the federal judicial hierarchy—to the US Appeals Court. (Given the small number of justices in our data in the US Supreme Court, our estimate is noisier and statistically indistinguishable from the probabilities for judges in the US District and Appeals courts.) In business, as firm size—a proxy for being part of the core—based on total assets (Panel E) and the number of employees (Panel F) increases, the likelihood of a firm demonstrating elite prescience declines precipitously (Table S7).

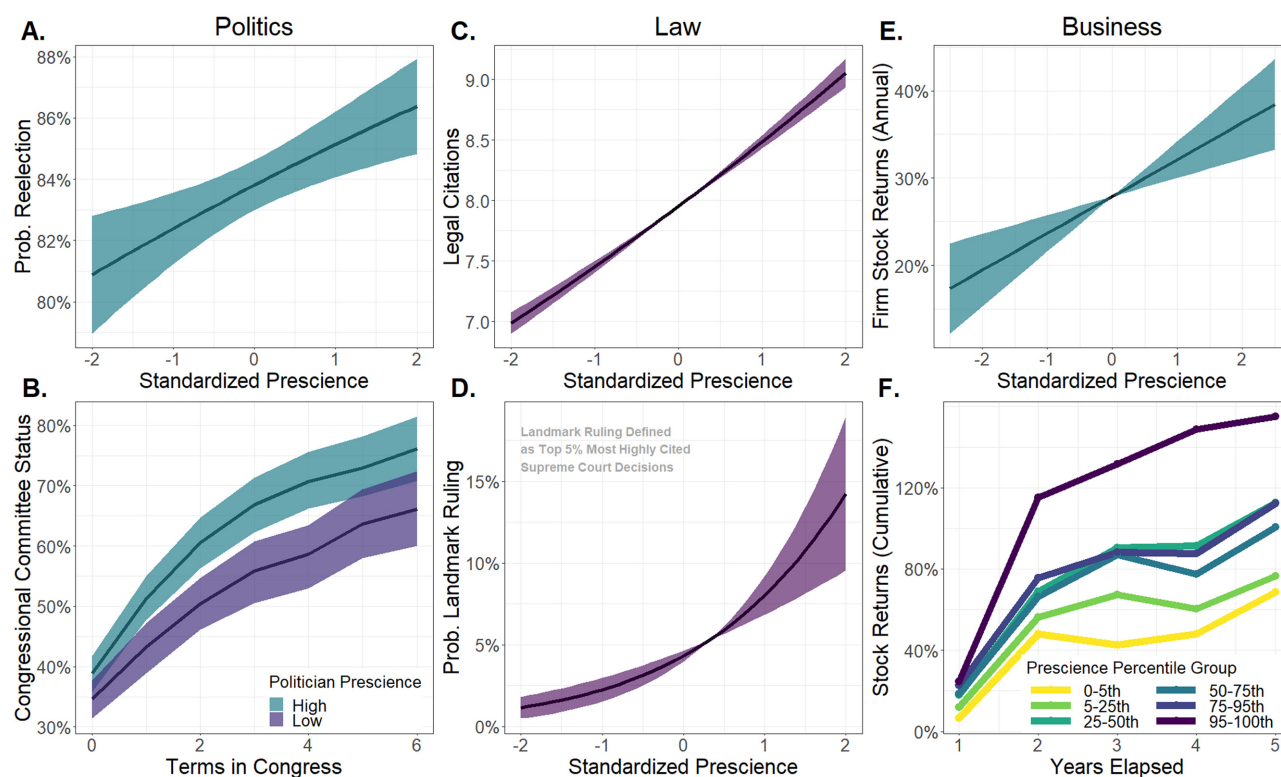


Fig. 4. Prescience predicts success (A) prescience predicts political reelection. Marginal effects plot from panel linear probability models of political reelection on standardized prescience; politician-term unit of analysis, with political party $\times$ congressional term fixed effects ($\beta = 0.00860$, $P < 0.05$). (B) Prescience predicts congressional committee status. Mean committee status for the top tercile (high prescience) and bottom tercile (low prescience) Congressional term with 10k politician bootstrapped SEs. Congressional committee status defined by the committee transfer ratio (Supplementary Material Appendix). Please see Supplementary Material for panel regressions with fixed effects and other controls ($\beta = 0.0155$, $P < 0.001$). (C) Prescience predicts highly cited court decisions. Marginal effects plot of linear regression model of log total citations on standardized prescience; judicial decision unit of analysis, with judge, court, and year fixed effects ($\beta = 0.0693$, $P < 0.001$). (D) Prescience predicts landmark Supreme Court decisions. Marginal effects plot of linear probability models of landmark decisions on standardized prescience. Landmark decision is defined as the top 5% most highly cited US Supreme Court decisions by year, with the sample restricted to US Supreme Court decisions. Models include judge and year fixed effects ($\beta = 0.687$, $P < 0.001$). (E) Prescience predicts firm stock returns. Marginal effects plot of linear regression models of yearly stock returns from 2012–2015 on 2011 standardized prescience; North American Industry Classification System (NAICS) three-digit industry fixed effects ($\beta = 0.0422$, $P < 0.001$). (F) Prescience predicts elite firm performance. Total stock returns since 2012 by prescience quartile and year (with top and bottom 5%). The y-axis shifts from annual stock returns in panel (E) to cumulative stock returns in panel (F).

Discussion

The ability to systematically identify the linguistic markers of prescient ideas enables us to shed light on longstanding questions about the origins of transformative ideas. Popular intuitions often suggest that prescient ideas emerge from powerful incumbents. In law, for example, higher court judges are typically thought to produce prescient rulings, which guide the subsequent judgments of lower-court judges (18). Similarly, in politics, established legislators who lead the most central committees are frequently identified as the most prescient (28). In business, by contrast, theories of disruptive innovation implicitly assume that the prescient ideas underpinning revolutionary products and business models arise from new entrants to an industry rather than from entrenched incumbents (3, 29). Because we have heretofore lacked a systematic way of quantifying prescience, such intuitions have been mostly informed by anecdotes and case studies. In contrast, our method reveals that, across a diverse array of domains, prescient ideas emanate from the periphery rather than the center.

Our results also suggest that prescience may be as much a property of contexts as of individuals. Indeed, in the legal domain, the same individual becomes less prescient as she moves up the status hierarchy to upper-level courts. Those in search of break-

through ideas should therefore look beyond the usual suspects who are ensconced in well-trodden places and instead focus attention on the unconventional ideas brewing in the outskirts of a domain.

A precondition for an idea to become prescient is that it is contextually novel, yet most novel ideas fail to gain traction (19–21). Thus, it remains unclear from our study whether a strategy of actively pursuing contextual novelty would lead to long-term success for an actor. We leave to future research the task of unpacking the conditions under which contextual novelty translates into prescience and eventual acclaim in a given field.

Although our empirical investigation focuses on three specific contexts, the method we introduce can detect prescient ideas in other domains. A natural extension, for example, would be the domain of science, where our empirical strategy could be readily applied to the text of academic papers and patents. Indeed, prior work in the “science of science” tradition has shown that disruptive innovation, as defined by citation networks, tends to originate in small and peripheral scientific teams rather than large and core ones (3, 6, 30). It remains to be explored the conditions under which prescient ideas, as determined by our measurement approach, translate into disruptive innovation as measured by citations to the tangible artifacts of scientific production. By focusing

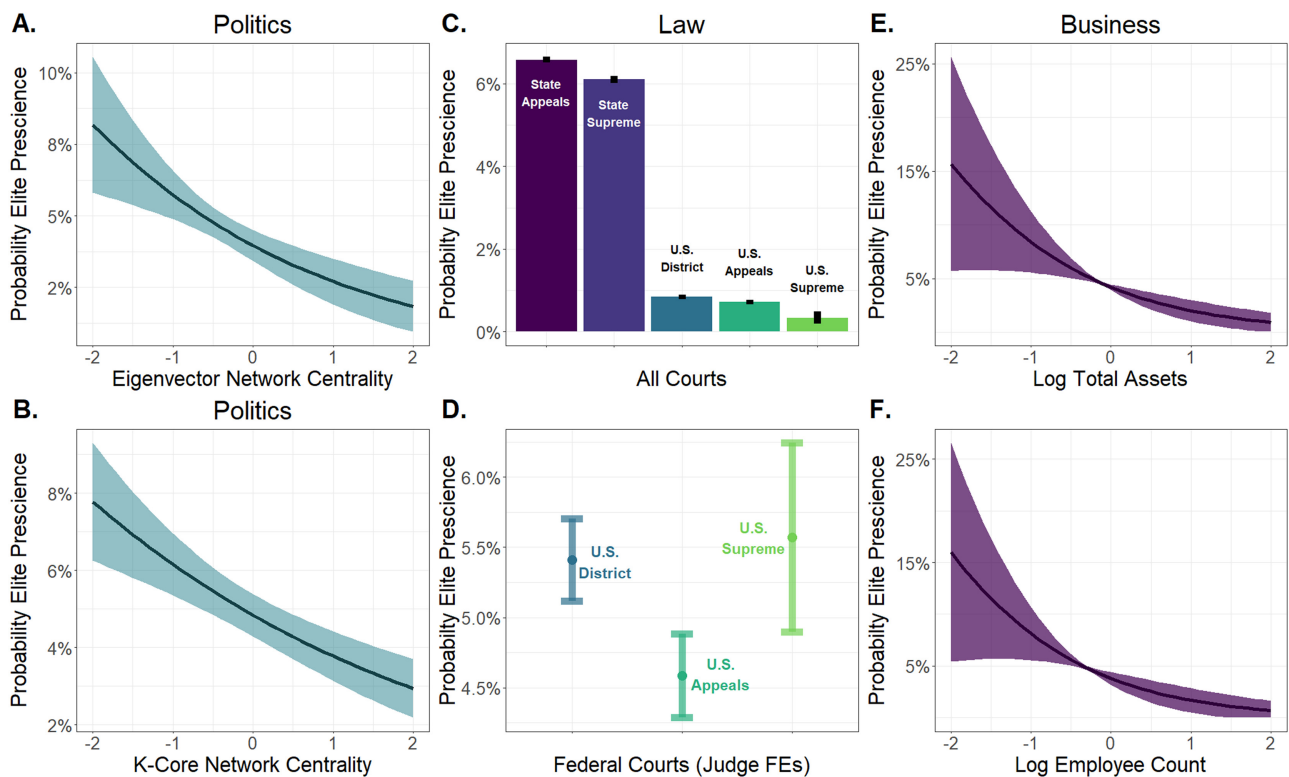


Fig. 5. Highly prescient ideas come from the periphery. Marginal effects plots regressing the probability of having elite prescience (top 5% in standardized prescience) on alternatives measures of peripheral positions using logistic regression and 95% confidence intervals. All regressions include controls for the log number of sentences and are restricted to observations with at least 50 sentences given increased variance in prescience with small sample size. (A and B) Highly prescient politicians come from peripheral network positions. Politician network defined using bill cosponsorship data (15). Network periphery measured by standardized eigenvector centrality ($\beta = -0.292$, $P < 0.001$) and standardized k-core centrality ($\beta = -0.277$, $P < 0.001$) with additional centrality measures in the Supplementary Material. (C and D) Highly prescient court decisions come from the lower courts. Panel (C) depicts the probability of a prescient decision using both state and federal courts and year fixed effects. Panel (D) adds judge fixed effects and restricts the sample to federal decisions (for which we have judge disambiguated decisions). (E and F) Highly prescient ideas come from small firms. NAICS two-digit industry fixed effects. Firm size measured by standardized total assets ($\beta = -0.207$, $P < 0.01$) and the standardized number of employees ($\beta = -0.309$, $P < 0.05$).

attention on and providing a novel means to quantify prescient ideas, we aim to broaden scholarly exploration from a narrow fixation on outcomes such as citations to the larger process by which prescient ideas that change the world emerge.

Materials and Methods

To extract prescience from conversational text data, we build upon a recent innovation in natural language processing and deep learning: BERT (14). BERT is a generalized language model, meaning it learns the syntax and semantic meaning of language, which can then be applied to a litany of downstream tasks like machine translation and named entity recognition. Underlying BERT is layers of transformer blocks. The transformer architecture diverges from previous language modeling approaches by replacing recurrence and convolutions with attention mechanisms (31). Doing so allows the entire sentence to be propagated through the model simultaneously, significantly speeding up parallelization. BERT stacks dozens of transformer blocks, encompassing hundreds of millions of parameters. As a result, BERT learns syntax relationships, semantic meanings, co-references, and even encodes entire syntax trees (32, 33). Because BERT is computationally intensive—often requiring several weeks of time on dedicated cloud tensor processing units (TPUs) to train on a new corpus—researchers typically begin with the pretrained model provided

by Google (where BERT was developed) and fine-tune this model to their own corpora. Through the fine-tuning process, the general meanings learned by BERT can be contextualized to the researchers' specific domains of interest (34).

Traditional language models process sentences left-to-right, one word at a time, and estimate the conditional likelihood of a word: $(w_i | w_0, w_1, \dots, w_{i-1})$. BERT instead favors bidirectionality—that is, it attends to both the left and right contexts simultaneously (14). To circumvent the unidirectional constraint, BERT is trained (in part) using a masked language model (MLM) task: 15% of the words in a sentence are randomly masked, and the model is tasked to predict the masked tokens.

Sentence (s): Earnings are up this quarter.
Masked s: Earnings are (MASK) this quarter.

The MLM objective differs from “true” language models in that the likelihood of the model generating a sentence is undefined. As a proxy, we use the model's ability to solve the MLM for each word in the sequence, leaving all other words unmasked. Here, the likelihood of a word is conditional on both the left and right contexts: $(w_i | w_0, w_1, \dots, w_{i-1}, w_{i+1}, w_{i+2}, \dots)$. While most researchers take the generalized language model features of BERT and add an additional layer on top to solve a downstream task, we instead directly use BERT's probabilistic modeling of language via MLM to quantify prescience.

Specifically, we task the model to minimize the cross-entropy loss between the predicted and the actual word. Let $\vec{y}_i$ represent a location in a vector of length N , where N refers to the number of words in the corpus, for word i . This one-hot encoded vector takes a value of one at the index of the masked token and zero otherwise. $\vec{\hat{y}}_i$, also of length N , is the vector predicted token likelihood obtained through a softmax activation layer predicting the masked token by the BERT model. Model accuracy is evaluated using cross-entropy loss, $-\vec{y}_i \cdot \log(\vec{\hat{y}}_i)$. To obtain word-level perplexity, PP_i , which is the inverse-likelihood of the model generating the word, we exponentiate the cross-entropy loss (Equation 1).

$$P_i = \exp(-\vec{y}_i \cdot \log(\vec{\hat{y}}_i)). \quad (1)$$

Words that are trivially predicted by the model—such as stop words and punctuation—when used in expected contexts have perplexities of ~ 1 , meaning that the model predicts them with close to 100% accuracy. Conversely, words and phrases that are highly unusual or unexpected have higher perplexity scores. We take the product of these perplexities and normalize by the n th root to account for sentence length (Equation 2).

$$CN(s) = \left(\prod_{i=0}^N PP_i \right)^{\frac{1}{N}}. \quad (2)$$

We refer to this term as CN instead of sentence-level perplexity for two reasons. First, given we use this measure to assess the extent to which ideas rethink the contextual assumptions in a domain, terming it contextual novelty aligns our empirical measure with our theoretical quantity of interest. Second, because BERT models bidirectionally, the perplexity of the sentence is technically undefined, and terming this sentence-level perplexity would be inconsistent with prior work (35).

To measure *prescience*, we rely on the intuition that prescient ideas depart from prevailing ideas at the time of introduction but become commonplace in the future (3, 36, 37). For example, Gerow et al. (23) identify highly influential scholarly publications by studying how academic discourse shifts after their publication. Thus, rather than fine-tuning BERT to our entire corpus, we split the corpus into time periods and fine-tune separate BERT models on each split of the corpus. This approach allows us to examine how the contextual novelty of a sentence changes over time.

Our method requires two BERT models trained on a corpus split into two periods, current (t_0) and future (t_1), and two BERT models that map documents to contextual novelties $CN_{t_0}(s)$, $CN_{t_1}(s)$. For a document from the current period, we define *prescience*, P , as the percentage reduction in contextual novelty between the current and future models.

$$P(s) = \frac{CN_{t_0}(s) - CN_{t_1}(s)}{CN_{t_0}(s)}. \quad (3)$$

Both *contextual novelty* and *prescience* are defined at the sentence level. To transition from the sentence level to the relevant unit of analysis in a given domain, we simply take the mean value of these variables over the unit of aggregation. For our corpus of the US State and Federal judicial decisions, for example, we aggregate at the unit of the judicial decision. We define truly prescient ideas (as manifested in individuals, organizations, or documents) as those in the top 5% of the distribution of prescience. We find that our measure is noisier with short sentences, as there is less context for BERT to use when making predictions. To reduce noise, we restrict to sentences with at least 10 tokens and <100 tokens (to catch errors in the sentence parser) before aggregating to

mean prescience. Researchers replicating this methodology may consider using a higher minimum token count, such as 15 or 20 tokens, to further reduce noise.

Data, fine-tuning, and measuring prescience

Training BERT from scratch is prohibitively expensive, taking weeks on a cloud TPU. Instead, Google has provided a pretrained model—trained on the BookCorpus (800 million tokens) and the English Wikipedia (2.5 billion tokens) available at <https://github.com/google-research/bert>—that researchers can fine-tune on their own corpora to learn context-specific idiosyncrasies. We fine-tune using BERT-base uncased (12-layer, 768-hidden, 12-head, 110 million parameter model) by repeating the MLM task and next-sentence prediction task on our corpora. We filter out sentences with <10 tokens to reduce noise and sentences >100 tokens given that they likely represent errors in the sentence parser. (An even higher minimum token count will greatly reduce noise in computing contextual novelty and prescience.) For fine-tuning, we use a max sequence length of 128, a batch size of 64, and fine-tune for $\sim 400,000$ steps. The only preprocessing of text prior to fine-tuning is sentence tokenization and appending (CLS) and (SEP) tokens to the start and end of each sentence, respectively. BERT uses WordPiece tokenization, which converts unrecognized tokens into subtokens [e.g., tokenizing onboarding into (onboard, ing)], so no out-of-vocabulary words are dropped from the analysis.

Below, we describe the three data sets in greater detail and explain the specific text preprocessing, fine-tuning, and approach to computing prescience we followed in each setting. For a full description, please refer to the Supplementary Material.

Politics

To identify prescient politicians, we use transcripts from the US House of Representatives and the US Senate from the bound and daily editions of the US Congressional Record from the 43rd to 114th Congresses (1873–2017). We use floor speeches, as opposed to other political texts such as press releases or campaign speeches, as these texts (1) maintain a consistent format and (2) offer a complete population of political discourse. We use the data set provided by Gentzkow, Shapiro, and Taddy (2019), who remove procedural language and parse the text from each congressional session into speeches attributable to congress people (38). (An even higher minimum token count will greatly reduce noise in computing contextual novelty and prescience.) For data reliability reasons [e.g., temporal variation in OCR accuracy], we begin our analyses with the 87th Congress, whose members took office in 1961 resulting in 4.9 million unique speech events. We focus on this later subset of the corpus because it contains fewer OCR errors, which can potentially bias our measure of prescience. We obtain biographical data on politicians from the Congressional Biographical Directory, GovTrack, and Congress.gov and collect committee membership (39, 40) and bill cosponsorship data (41).

We segment the corpus into 4-y increments, corresponding to presidential terms, resulting in 15 buckets—1961–1964, 1965–1968, ... 2007–2010—and fine-tune separate BERT models for each increment. We compartmentalize corpus by presidential terms given that these are natural breakpoints in the trajectory of political discourse. To compute *prescience*, we define the current period as the BERT model trained using the year the sentence was spoken, M_0 . A more difficult choice is in selecting the future period. Choosing a proximate model quantifies short-term evolution in discourse, while a model trained on text further in time from the

focal sentence quantifies longer-term prescience. To balance this trade-off, we select two future period models—the immediately subsequent model M_1 and the model after that M_2 —and take the arithmetic mean of prescience between these two prescience calculations: $[P(M_0, M_1) + P(M_0, M_2)]/2$. We define the current period model as the BERT model trained using the year the sentence was spoken. So, for example, for a sentence spoken in 1965, we use the BERT model trained on sentences from 1965 to 1968 as M_0 and the 1969–1972 and 1973–1976 models as the future models (3).

Law

Our data of US State and Federal cases comes from the Caselaw Access Project, which digitized and processed >8 million state and federal judicial verdicts, stretching back to the 1640s. These data provide the complete population of legal discourse and are the main forum of judicial interpretation and legal precedent. As with our data on political discourse, we found significantly more OCR errors prior to 1960. Thus, to both align these data with data on politicians and to minimize OCR errors, we restrict our analysis to cases beginning in 1960. Our resulting sample includes 4.2 million cases. We obtain biographical data on federal judges, including their court tenure, gender, and prior judicial service from the Federal Judicial Center. Data on case citations come from the citation graph provided by the Caselaw Access Project, which extracts citations from the in-line text of court decisions. We remove in-line citations using LexNLP, a python package specifically designed to parse legal text, and sentence tokenize using this package as well.

We compartmentalize the corpus into 5-y intervals—1960–1964, 1965–1969, . . . , 2005–2009—and fine-tune separate BERT models on each interval. As with the politics data set, we define the current period model as the BERT model trained using the focal year and compute mean prescience using the subsequent two models to strike a balance between assessing shorter-term versus longer-term prescience.

Business

To study prescient ideas in business, we collect a corpus of QECs from seekingAlpha, a content service for financial markets. Our dataset includes 108,334 QECs (414 million tokens) from publicly traded firms, predominately headquartered in the United States from 2006 to 2016. We restrict our analyses to the Q and A section of the call, removing the prepared remarks by the company and filtering out statements by analysts and the operator. We restrict to the Q and A portion because prepared statements have a very different style of discourse than the Q and A section, which may complicate fine-tuning on a smaller corpus. The back-and-forth discussion in QECs, unlike alternative forms of corporate text such as annual reports or press releases, offers intimate insights into the potentially prescient ideas of firms. We disambiguate company names in QECs and fuzzy match them to Compustat to obtain firm characteristics and performance outcomes. We identify 5,847 firms that have corresponding links to Compustat gvkeys. We then match gvkeys with Permno to link to the CRSP database, thereby allowing us to collect data on daily stock returns.

Unlike our datasets of political speeches and legal decisions, this corpus is comparatively small. We lack a significant number of transcripts for speeches between 2006 and 2011 and the text in the earlier speeches is heavily influenced by the 2007/2008 financial crisis. As a result, we restrict our analysis to QECs from 2011 onward. Given that these data span a relatively small number of years, we split the corpus on an annual basis and fine-tune a separate BERT model for each year in our data. We select 2011 as the focal year and select both 2015 and 2016, the last 2 y in our data,

as the future comparison periods when computing prescience. Because the unit of analysis is the firm, we aggregate all QECs for each firm in 2011 by computing the mean prescience across calls in a given year to create a firm-level measure of prescience.

Validating prescience

In Figure 4, we demonstrate how prescient actors reap rewards. For more details on these analyses and their associated regression tables, please refer to the Supplementary Material.

Politics

We begin by considering the extent to which prescience appears to be rewarded for politicians in our data. We aggregate political prescience at the congressional-term unit of analysis and select two measures of political success. The first, political reelection, is widely considered the primary motivator of incumbent politicians (42). The second variable, *committee status*, measures political success within the legislative chamber itself. We use the transfer ratio defined by Bullock and Sprague in 1969, which is computed for each committee (43). We estimate linear regressions of reelection and committee status on average prescience (Table S1).

Law

We turn next to considering whether prescience is associated with success in the legal domain. We define success as the number of legal citations that accrue to the case. We restrict our analyses of citations to federal cases because doing so enables us to include judge-level controls, such as judge fixed effects, which are unavailable for state-level decisions. We also restrict to cases with at least 50 sentences to reduce variance in prescience. This restriction drops <1% of decisions.

The first dependent variable of interest, *log citations*, is the natural log of the number of citations a case receives plus 1. We log transform the variable to reduce skew. A more stringent test of our method is whether it can identify landmark decisions—judicial rulings that significantly alter existing interpretations of existing law. Most landmark decisions come from the US Supreme Court, given they have absolute authority to set precedents and determine law. We define a US Supreme Court ruling as a *landmark decision* if it is above the 95th percentile in citations for the court. We estimate *log citations* using linear regression and *landmark decision* using logistic regression (Table S2).

Business

To validate our measure of prescience in the business context, we began by testing whether our measure of prescience can predict future stock returns. We measure stock returns using daily returns data: the percentage change in value accrued to the stockholder on a given day. We report the model looking at 3-y returns from 2012 to 2014 in Figure 4, panel E, although our findings are robust to all windows (Table S5). We use three-digit North American Industry Classification System (NAICS) fixed effects in all regressions to adjust for industry-level heterogeneity (Tables S3 and S4).

Prescience arises from the periphery of a field

Next, we turn to the core claim of the paper: Prescient ideas are more likely to emerge from the periphery than from the center. There are many different ways to operationalize core versus peripheral positions—for example, based on an actor's position in the network of actors, relative size, and social status (e.g., based on demographic traits such as gender or elite credentialing). We

use each of these different approaches in our three empirical settings. The choice of how to operationalize core versus peripheral positions is based on our understanding of each empirical context and the nature of the data that were available to us. The following section describes in detail the data and analyses used to create Figure 5.

We define highly prescient politicians, judicial decisions, and firms as the top 5% in prescience during their relevant measurement period. For politicians, this means the top 5% in their current legislative term. We select the 95th percentile and above as prior work predicting scientific impact uses this threshold (3, 20). When predicting *highly prescient*, we control for the number of sentences for each actor, because those with relatively few utterances are disproportionately likely to exhibit high levels of prescience simply because our measure of prescience is noisier with fewer data. We do so using *log sentence count*.

Politics

We define core and peripheral positions in the legislative body via the politician's network position, as prior research demonstrates that well-connected politicians can influence their peers and governmental policy (41). To reconstruct the network of relationships between politicians, we use data on bill cosponsorships (44–46) and construct separate bill cosponsorship networks for each congressional term. We use a variety of network statistics to measure how peripheral a politician is in the cosponsorship network, including *degree centrality*, *eigenvector centrality*, *closeness centrality*, and *k-core centrality*. Because network centrality measures vary over time and between chambers, we standardize each network measure at the congressional term $\times$ chamber level. We regress *highly prescient politician* on our suite of network covariates using logistic regression (Table S5).

Law

In our legal setting, the US Supreme Court is most firmly placed in the core because it has statutory authority over all other courts. Recognizing that jurisdictional differences render it difficult to make direct comparisons, we assume that state courts are more peripheral than federal courts because the fragmented body of state laws is consolidated by federal rulings but not vice versa. Thus, our judicial hierarchy goes from most peripheral to most core in the following order: State Appeals, State Supreme, US District, US Appeals, and US Supreme. We define *highly prescient decision* as the top 5% most prescient decisions in a given year across all courts. To test the idea that prescient judicial decisions come from the structural periphery of the court system, we regress *highly prescient decision* on separate indicator variables for each court using logistic regression (Table S6).

Business

Whereas the periphery in the legal system is defined by hierarchy, we define the periphery in business based on firm size. Large firms command significant market power and use their power to lobby congress for favorable regulations, set industry standards, and influence prices. By virtue of their size, large firms also occupy more central positions in the network between firms (47). We define firm size using *log total assets*, the total amount of economic value held by the firm, and its number of employees, operationalized as *log employee count*. We define a *highly prescient firm* as one at or above the 95th percentile in prescience and estimate using logistic regression (Table S7).

Acknowledgments

We thank Yixi Chen for her invaluable research assistance in collecting and processing the data. We also thank Dan McFarland and Dashun Wang for their helpful comments on prior drafts. The usual disclaimer applies.

Supplementary Material

Supplementary material is available at [PNAS Nexus](https://pnas.nexus) online.

Authors' Contributions

- Conceptualization: P.V., A.G., and S.B.V.
- Data curation: P.V., A.G., and S.B.V.
- Methodology: P.V.
- Investigation: P.V., A.G., and S.B.V.
- Visualization: P.V.
- Writing—original draft: P.V., A.G., and S.B.V.
- Writing—review and editing: P.V., A.G., and S.B.V.

Preprints

The manuscript was posted on a preprint: <https://osf.io/preprints/socarxiv/j24pw/>.

Data Availability

Some data used in the study are obtained from proprietary databases such as the Center for Research in Security Prices and Compustat and cannot be made available. The raw textual data, such as political transcripts and legal documents, are made available in other repositories cited in the manuscript. The code needed to compute prescience is provided on the author's GitHub (<https://github.com/pvicinanza>) and fine-tuned BERT models are available upon request.

References

1. Isaacson W. 2011. *Steve Jobs: the exclusive biography*. New York (NY): Simon & Schuster.
2. Roberts A. 2015. *Napoleon: a life*. Preprint ed. New York (NY): Penguin Publishing Books.
3. Wu L, Wang D, Evans JA. 2019. Large teams develop and small teams disrupt science and technology. *Nature*. 566(7744):378–382.
4. Singh J, Fleming L. 2010. Lone inventors as sources of breakthroughs: myth or reality? *Manag Sci*. 56(1):41–56.
5. Hofstra B, et al. 2020. The diversity–innovation paradox in science. *Proc Natl Acad Sci*. 117(17):9284–9291.
6. Wuchty S, Jones BF, Uzzi B. 2007. The increasing dominance of teams in production of knowledge. *Science*. 316(5827):1036–1039.
7. Ma Y, Uzzi B. 2018. Scientific prize network predicts who pushes the boundaries of science. *Proc Natl Acad Sci*. 115(50):12608–12615.
8. Curtis JN. 2017. Remembering racial progress, forgetting white resistance: the death of Mississippi Senator John C. Stennis and the consolidation of the colorblind consensus. *Hist Mem*. 29(1):134–160.
9. Kuhn TS. 2012. *The structure of scientific revolutions: 50th anniversary edition*. Chicago (IL): University of Chicago Press.

10. Garg N, Schiebinger L, Jurafsky D, Zou J. 2018. Word embeddings quantify 100 years of gender and ethnic stereotypes. *Proc Natl Acad Sci.* 115(16):E3635–E3644.
11. Grand G, Blank IA, Pereira F, Fedorenko E. 2022. Semantic projection recovers rich human knowledge of multiple object features from word embeddings. *Nat Hum Behav.* 6:975–987.
12. Rodriguez PL, Spirling A. 2022. Word embeddings: what works, what doesn't, and how to tell the difference for applied research. *J Polit.* 84(1):101–115.
13. Kozlowski AC, Taddy M, Evans JA. 2022. The geometry of culture: analyzing the meanings of class through word embeddings. *Am Sociol Rev.* 84(5):905–949.
14. Devlin J, Chang M-W, Lee K, Toutanova K. 2019. BERT: pre-training of deep bidirectional transformers for language understanding. In: Burstein J, Doran C, Solorio T, editors. *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1.* Minneapolis (MN): Association for Computational Linguistics. p. 4171–4186.
15. Merton RK. 1973. *The sociology of science: theoretical and empirical investigations.* Chicago (IL): University of Chicago Press.
16. Borgatti SP, Everett MG. 2000. Models of core/periphery structures. *Soc Netw.* 21(4):375–395.
17. Cattani G, Ferriani S. 2008. A core/periphery perspective on individual creative performance: social networks and cinematic achievements in the Hollywood film industry. *Organ Sci.* 19(6):824–844.
18. Schultz DA, Smith CE. 1996. *The jurisprudential vision of Justice Antonin Scalia.* Lanham (MD): Rowman & Littlefield.
19. Fleming L. 2001. Recombinant uncertainty in technological search. *Manag Sci.* 47(1):117–132.
20. Uzzi B, Mukherjee S, Stringer M, Jones B. 2013. Atypical combinations and scientific impact. *Science.* 342(6157):468–472.
21. Foster JG, Rzhetsky A, Evans JA. 2015. Tradition and innovation in scientists' research strategies. *Am Sociol Rev* 80(5):875–908.
22. Coppola L, DeMarco N. 2012. Civil RICO: how ambiguity allowed the racketeer influenced and corrupt organizations act to expand beyond its intended purpose. *N Engl J Crim Civ Confin.* 38(2):241–256.
23. Gerow A, Hu Y, Boyd-Graber J, Blei DM, Evans JA. 2018. Measuring discursive influence across scholarship. *Proc Natl Acad Sci.* 115(13):3308–3313.
24. Arts S, Hou J, Gomez JC. 2021. Natural language processing to identify the creation and impact of new technologies in patent text: code, data, and new measures. *Res Policy.* 50(2):104144.
25. Mann R. 1992. *Legacy to power: Senator Russell Long of Louisiana.* New York (NY): Paragon House.
26. Theriault SM. 2013. *The Gingrich Senators: the roots of partisan warfare in Congress.* New York (NY): Oxford University Press.
27. Knoke D, Pappi FU, Broadbent J, Tsujinaka Y. 1996. *Comparing policy networks: labor politics in the US, Germany, and Japan.* Cambridge (UK): Cambridge University Press.
28. Diller D, Stefani S. 2019. *Richard G. Lugar: Indiana's visionary statesman.* Bloomington (IN): Indiana University Press.
29. Christensen C, McDonald RM, Altman EJ, Palmer J. 2017. *Disruptive innovation: intellectual history and future paths.* Harvard Business School. Working Paper No. 17-057.
30. Jones BF, Wuchty S, Uzzi B. 2008. Multi-university research teams: shifting impact, geography, and stratification in science. *Science.* 322(5905):1259–1262.
31. Vaswani A et al., 2017. Attention is all you need. In: von Luxburg U, Guyon I, Bengio S, Wallach H, Fergus R, editors. *31st Conference on Neural Information Processing Systems.* Long Beach (CA): NeurIPS. p.6000–6010.
32. Hewitt J, Manning CD. 2019. A structural probe for finding syntax in word representations. In: Burstein J, Doran C, Solorio T, editors. *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics.* Minneapolis (MN): Association for Computational Linguistics. p.4129–4138.
33. Clark K, Khandelwal U, Levy O, Manning CD. 2019. What does BERT look at? an analysis of BERT's attention. In: Linzen T, Chrupala G, Belinkov Y, Hupkes D, editors. *Proceedings of the 2019 ACL workshop BlackboxNLP: analyzing and interpreting neural networks for NLP.* p. 276–286.
34. Bengio Y. 2012. Deep learning of representations for unsupervised and transfer learning. In: Guyon I, Dror G, Lemaire V, Taylor G, Silver D, editors. *Proceedings of ICML workshop on unsupervised and transfer learning.* Bellevue (WA): PMLR. p. 17–37.
35. Rosenfeld R. 2000. Two decades of statistical language modeling: where do we go from here? *Proc IEEE.* 88(8): 1270–1278.
36. Kelly B, Papanikolaou D, Seru A, Taddy M. 2021. Measuring technological innovation over the long run. *Am Econ Rev: Insights.* 3(3):303–320.
37. Funk RJ, Owen-Smith J. 2016. A dynamic network measure of technological change. *Manag Sci.* 63(3):791–817.
38. Gentzkow M, Shapiro JM, Taddy M. 2019. Measuring group differences in high-dimensional choices: method and application to congressional speech. *Econometrica.* 87(4): 1307–1340.
39. Stewart C, III. 1992. Committee hierarchies in the Modernizing House, 1875-1947. *Am J Pol Sci.* 36(4):835–856.
40. Stewart C, III, Woon J. 2017. *Congressional Committee Assignments, 103rd to 114th Congresses, 1993–2017.*
41. Fowler JH. 2006. Connecting the Congress: a study of cosponsorship networks. *Polit Anal.* 14(4):456–487.
42. Schlesinger JA. 1966. *Ambition and politics: political careers in the United States.* Chicago (IL): Rand McNally and Co.
43. Bullock C, Sprague J. 1969. A research note on the committee reassignments of Southern Democratic Congressmen. *J Polit.* 31(2):493–512.
44. Fowler JH. 2006. Legislative cosponsorship networks in the US House and Senate. *Soc Netw.* 28(4):454–465.
45. Brinke LT, Liu CC, Keltner D, Srivastava SB. 2016. Virtues, vices, and political influence in the US Senate. *Psychol Sci.* 27(1): 85–93.
46. Liu CC, Srivastava SB. 2019. Efficacy or rigidity? power, influence, and social learning in the US Senate, 1973–2005. *Acad Manag Discov.* 5(3):251–265.
47. Allen MP. 1974. The structure of interorganizational elite cooperation: interlocking corporate directorates. *Am Sociol Rev.* 39(3):393–406.