



Artificial intelligence and scientific discovery: a model of prioritized search[☆]

Ajay Agrawal^a, John McHale^b, Alexander Oettl^{c,*}

^a University of Toronto and NBER, Canada

^b University of Galway, J.E. Cairnes School of Business and Economics, Ireland

^c Georgia Institute of Technology and NBER, United States of America

ARTICLE INFO

Keywords:

Artificial intelligence
Scientific discovery
Scientific Search
Innovation
Theory

ABSTRACT

We model a key step in the innovation process, hypothesis generation, as the making of predictions over a vast combinatorial space. Traditionally, scientists and innovators use theory or intuition to guide their search. Increasingly, however, they use artificial intelligence (AI) instead. We model innovation as resulting from sequential search over a combinatorial design space, where the prioritization of costly tests is achieved using a predictive model. The predictive model's ranked output is represented as a hazard function. Discrete survival analysis is used to obtain the main innovation outcomes of interest – the probability of innovation, expected search duration, and expected profit. We describe conditions under which shifting from the traditional method of hypothesis generation, using theory or intuition, to instead using AI that generates higher fidelity predictions, results in a higher likelihood of successful innovation, shorter search durations, and higher expected profits. We then explore the complementarity between hypothesis generation and hypothesis testing; potential gains from AI may not be realized without significant investment in testing capacity. We discuss the policy implications.

1. Introduction

The recent explosion in interest in AI has predominantly focused on the dramatic improvements in large language models (LLMs), reflected in applications like ChatGPT. Public discussion has highlighted both the uncertain costs and benefits. Understandably, the potential costs have received most attention, including the adverse impacts on labor markets due to task automation, the lowered cost of producing misinformation, the dangers of AI-facilitated surveillance, and the damage from biases in automated decision processes in areas ranging from the provision of credit to the granting of bail (see, e.g. [Acemoglu and Johnson, 2023](#)).

Meanwhile, a less public – but influential – view has developed that AI's truly transformative potential lies in its ability to change the process of scientific discovery and innovation ([Hassabis, 2022](#); [Krenn et al., 2022](#)). The latter perspective is predicated on AI's potential to provide prediction tools to guide search over vast and complex combinatorial spaces ([Bianchini et al., 2022](#)). So far, the most celebrated use of AI for scientific discovery is AlphaFold, which predicts the 3D shape of proteins based on their amino-acid sequences ([Jumper et al., 2021](#); [Cavalli, 2023](#)). The potential for AI-driven scientific discovery extends beyond this, with growing applications in fields ranging from medicine to materials science.¹ However, existing models of the innovation do not

[☆] We thank Jefferson Galetti, Jason Harold, Akhil Sasidharan, and Anil Yadav for helpful feedback. We also thank the editor and three anonymous referees for their helpful advice. Of course, responsibility for all remaining deficiencies is our own. Finally, we gratefully acknowledge Science Foundation Ireland (Grant Number 17/SPR/5329) for financial support.

* Corresponding author.

E-mail address: alexander.oettl@scheller.gatech.edu (A. Oettl).

¹ See [Ramsundar et al. \(2019\)](#) for an accessible introduction to the use of deep learning in the life sciences. [Raghu and Schmidt \(2020\)](#) survey deep learning techniques used in scientific discovery.

incorporate the role of the prediction in the innovation process.

Our contribution in this paper is to present a tractable model of the innovation process designed to highlight the role of predictive models in prioritizing search over vast and complex combinatorial design spaces. AI can be viewed as a technology for discovery to the extent it provides a general method for developing predictive models over such spaces.² As such, AI is a tool for hypothesis generation that may be used as a substitute or a complement for traditional approaches, such as using theory or intuition for generating hypotheses (predictions).

Notwithstanding the rapid pace of technological advance, some worry that we have yet to see significant evidence of improved productivity in the innovation process. This may turn out to be similar to what happened with previous general-purpose technologies (GPTs), where market failures slowed the advancement and evolution of the technology, limiting its short-term impact on productivity, market valuations, and income distribution (David, 1990; Bresnahan and Trajtenberg, 1995; Jovanovic and Rousseau, 2005; Bresnahan, 2010; Brynjolfsson et al., 2019, 2021).³ As a special kind of GPT, AI fits the category of an “invention of a method of invention” (Griliches, 1957; Cockburn et al., 2019; Bianchini et al., 2022) or “meta-GPT” (Romer, 2008; Agrawal et al., 2019). Relative to “regular” GPTs, a meta-GPT has the potential for outsized impacts to the extent it changes the knowledge production function of the economy (Rosenberg, 1998).⁴

The complexity of the new research frontiers provides one explanation for disappointing recent productivity growth despite exponentially increasing resources invested in R&D.⁵ Innovators lack the kinds of science-based predictive theories to search these spaces that underpinned past scientific advances, most notably the predictive tools provided by 19th and 20th century advances in theoretical physics. Many of today’s greatest innovation challenges are in the domain of biology – small molecule-protein binding, gene expression, protein structure prediction, protein design, etc.⁶ Scientists hope that AI-based predictive models will provide useful tools to guide search at these challenging frontiers.

We use a special case of Weitzman sequential search (Weitzman, 1979) to model the search process. In our baseline approach, a predictive model – possibly AI based – is used to rank potential designs by their probability of success as generated by that model. This ranking is then inverted to produce a discrete hazard function. The hazard function represents the information available to prioritize the sequential search. Armed with the hazard function, we then apply discrete survival analysis to identify the expected innovation outcomes of interest. Importantly, our model specifies a causal relationship between predictive

power and innovation output, search duration (amount of testing before finding a success), and profitability.

After setting up this machinery, we then examine the impact of improving the fidelity of predictions – for example, by shifting from traditional methods of hypothesis generation, using theory or intuition, to using artificial intelligence instead. We present comparative statics that illustrate how increasing the quality of hypothesis generation via enhanced prediction fidelity increases innovation output, reduces the amount of testing required to find a success, and thus enhances profitability.

Our baseline model assumes that discoveries are binary (e.g., drug discovery successful or not). However, in many domains heterogeneity in success outcomes is possible (e.g., new materials discovery that achieves varying levels of success on all properties of interest, including the cost of production). So, we extend our baseline model by relaxing the assumption of binary outcomes and show the conditions under which riskier hypotheses (high payoff/low probability potential combinations) are prioritized. We also extend our baseline model to consider the competition-related choice scientists face between pursuing discoveries for more crowded problems with big prizes versus less crowded problems with smaller prizes. In addition, we extend our baseline model to explore the scientist’s exploration-exploitation decision. Scientists can choose to test hypotheses that are less likely to be successful but more likely to increase future predictive power by generating data in areas of the search space that are otherwise sparsely populated.

Increasing the speed and fidelity of hypothesis generation highlights the increasing cost of the bottleneck imposed by hypothesis testing. For most of the paper, we treat testing as an exogenous variable cost. However, after developing a baseline model and a series of extensions, all focused on the first step of the innovation process, hypothesis generation, we turn our attention to the second step, hypothesis testing.

We endogenize the testing decision. Tests can be conducted in series or parallel. The upside of testing in parallel is speed. Multiple hypotheses can be tested at once. The downside of testing in parallel is the lost opportunity to stop a search when a success is found. With sequential testing, the scientist can stop testing when a success is found, potentially reducing the time and number of tests required to discover a success. So, with no penalty for longer searches, sequential searches will result in higher innovation output and higher expected profits.

For realism, we introduce a penalty for slowness by imposing time discounting. Then, we relax the constraint that all tests take the same length of time. Tests can be slow or fast. Scientists can invest to speed up their tests. We examine the interaction between parallel-vs-sequential and slow-vs-fast testing. We show that sequential testing benefits more from speed than parallel testing. Thus, there are two potential equilibria: slow-parallel testing and fast-sequential testing. Given the spillovers that result from science, there may be significant welfare gains from policies that shift science systems from slow-parallel to fast-sequential testing regimes. Perhaps unsurprisingly, governments have begun subsidizing research into “self-driving labs” that fully automate the discovery loop: an AI generates a hypothesis, a robot tests the hypothesis, the test generates an outcome, the AI learns from the outcome, the AI updates its hypothesis, the robot tests the new hypothesis, and the feedback loop

² In a recent paper, Ludwig and Mullainathan (2023) consider an alternative way that AI can be used for hypothesis generation. They consider a situation in which an AI can find patterns not accessible to a human observer. Thus, the AI can suggest hypotheses that are both interpretable and novel. In the settings we assume, the search landscapes can be highly “rugged,” especially in the biological sciences, and interpretability may be challenging. The focus is often on finding something that works – e.g., an amino-acid sequence that produces a protein with some desired function – and it may not be possible to understand (at least initially) why it works. However, these two types of hypothesis generation can be complementary, with the finding of something that works leading to the development of new theories of causal mechanism (Rosenberg, 1994; Mokyr, 2002; Krenn et al., 2022).

³ Hötte et al. (2022) examine the issues involved in the measurement of AI as a GPT.

⁴ Crafts (2021) identifies meta-GPTs with industrial revolutions, and speculates that AI could facilitate the “fourth industrial revolution.”

⁵ See, e.g., Bloom et al. (2020).

⁶ Chemistry (including materials science) provides an interesting intermediate case. While physics-based predictive theories are in principle available, the vast space of possible molecules makes it practically difficult to make predictions, even using modern simulation techniques, across the entire relevant space.

continues.⁷

Ultimately, this paper is anchored in what we have named the ‘Hassabis hypothesis’, inspired by Demis Hassabis, CEO and c-founder of DeepMind. This hypothesis contends that leveraging AI-based predictive models can significantly advance scientific discovery and innovation. Hassabis outlines specific criteria for an AI-friendly scientific problem: a broad search space, a defined objective for model training, and substantial data or simulation capabilities (Hassabis, 2022). We propose a fourth criterion: the relative ineffectiveness of current predictive models in exploring the design space. This hypothesis implies a wide array of challenges uniquely addressable through AI, as demonstrated by breakthroughs like AlphaFold.

We structure the rest of the paper as follows. In Section 2, we develop our baseline model and in Section 3, we apply the model to show how an improvement in the available predictive model could lead to “better, faster, cheaper” innovation. In Section 4, we consider various extensions to the baseline model and in Section 5, we extend it further to examine the idea of an autonomous discovery system. We conclude in Section 6 with brief discussions of policy implications and the main testable implication of the model.

2. Baseline model

We model innovation as prioritized costly search over a combinatorial design space. The combinatorial approach to thinking about innovation has a long history in economics (Usher, 1929; Schumpeter, 1939; Nelson and Winter, 1982; Weitzman, 1998; Fleming and Sorenson, 2004; Arthur, 2009; Clancy, 2018; Jones, 2021). Agrawal et al. (2019) develop a model of AI-aided search over a combinatorial *idea* space. Such a focus on combining ideas might be appropriate if AIs approach truly human-like general intelligence. However, the current uses of AI in discovery seem better described as search over typically fixed *design* spaces – the spaces of molecules, amino acid sequences, lines of software code, etc. – and we concentrate on modelling the use of narrow AI predictive tools in prioritizing search in such settings. Although the type of statistics-based AI we consider is narrow,⁸ the range of its potential innovation applications makes it a meta-GPT.

2.1. A model of hypothesis generation: Outline

We develop a model of innovation as prioritized costly search over a combinatorial design space. Prioritization is achieved through a predictive model, which could involve human or artificial intelligence or some combination of the two. Candidate designs (i.e., potential combinations) must undergo costly testing before deployment. There are therefore, two stages to the innovation process – predict and test.

As an example, consider the problem of finding a small molecule

drug that binds with a target protein to achieve some desired therapeutic effect. The search space for hypothesis generation is potentially vast – e.g., the space of all molecular compounds that can be combined from N molecular “elements” – and the innovator is uncertain as to the locations of the combinations that would be successful for a given target. For simplicity, we assume that for the N elements, each element can be part of a combination (value = 1) or not part of the combination (value = 0). Given N elements, there are therefore 2^N potential combinations (each described by a N dimensional vector) that could be searched.

We first briefly preview the solution strategy (Fig. 1). Given the combinations that comprise the design search space, we start with the probability generated by a predictive model for each element of the set of possible combinations.⁹ Under sequential search, the ranking of these combinations by probability of success is used to define a discrete hazard function. This allows tools from discrete survival analysis to capture the key expected economic outcomes from the innovation process. The result is a tractable model linking the predictive model to economic outcomes of interest. These outcomes are: the optimal *maximum* duration of a completed search; the expected innovation output (and value) from a completed search; the expected duration (and cost) of a completed search; and, putting the pieces together, the expected profit from a completed search.

We make a number of assumptions to show the logic of the model in a relatively simple setting. (We discuss how these assumptions could be relaxed in Section 4.) First, there is *no heterogeneity* in outcomes conditional on a combination being successful in testing – e.g., there is no differentiation between successful drugs. Second, there is *no competition* between innovators working on a given problem – e.g., innovators are not racing to discover a drug to treat a target disease. And third, there is *no exploration*, so that an innovator seeks simply to exploit a given predictive model to maximize the expected value from a given innovation search – e.g., a drug company is not engaging in exploratory search with a goal of improving the predictive model.

2.2. Optimal sequential search and the hazard function

To model hypothesis generation (search), we adopt a special case of the Weitzman (1979) sequential search model. For any combination i that represents a potential design, we denote the output from the predictive model, the probability of success, as q_i . We assume that the innovator has access to unbiased predictions in the sense that for any set of combinations that produce a predicted probability of success equal to q , the proportion of successes in this set is also equal to q . In developing the prediction model, we assume that the data on prior successes and failures is an IID sample from the design space of potential contributions. Our risk-neutral innovator achieves a fixed payoff v from a (first) success that we normalize to 1 without loss of generality. Testing a combination (or opening a “box” in Weitzman’s terminology) has a common cost, c , and can be viewed as a draw from a Bernoulli distribution given the expected probability of success for that combination.

The assumptions of common values and common costs across combinations is what allows us to rank combinations solely on the basis of their probabilities of success. The variation in Weitzman reservation

⁷ In April 2023, the Government of Canada awarded “one of the largest federal research grants in Canadian history [to] making chemistry research automatic.” The goal of this research is “to develop self-driving labs, in an effort to automate and vastly accelerate the pace of chemical and materials research.” (Toronto Star, April 28, 2023).

⁸ Notwithstanding the disappointments with rules-based AI, the recent successes of statistical approaches to AI – most notably deep learning and self-supervised learning – have renewed interest in the potential for AI as a technology for innovation. While some AI scholars see such statistical-learning-based approaches as the best route to artificial general intelligence (AGI), the focus has largely been on their use for more narrow prediction tasks that are part of a broader innovation process that retains a central guiding role for human intelligence.

⁹ We abstract from the details of any particular real world problem and simply assume that a predictive model is available that attaches a probability to each combination in the design space.

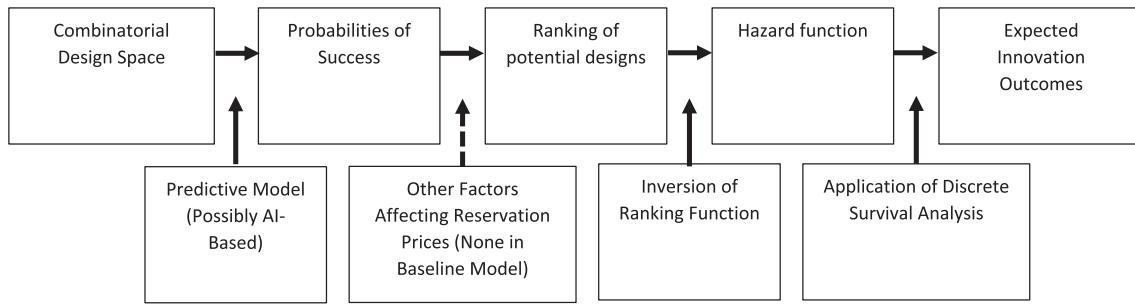


Fig. 1. Outline of a model of the innovation process with an exogenous testing cost. Note: The goal of the model is to link a predictive model over a combinatorial design space to the expected innovation outcomes. We assume the innovator engages in Weitzman sequential search. A test of a design is modelled as a Bernoulli trial. With no heterogeneity in the value of a success or the cost of a test in the baseline model, reservation prices (and thus the ranking of designs) are completely determined by the probabilities of success that are the outputs of the predictive model. Inverting the ranking based on these probabilities results in a discrete hazard function – i.e., a function that gives the probability of success conditional on reaching that point in the ranking in a sequential search. Discrete survival analysis is then used to derive the key expected innovation outcomes of interest.

prices is fully determined by variation in the probabilities of success.¹⁰ We can therefore think of the innovator as searching over a large set of Bernoulli distributions, one for each design in the search space.

We assume that a *positive* test is determinative in the sense that the false positive rate on the test is zero.¹¹ We initially assume there is no time discounting, though we relax that assumption in Section 5.¹²

The optimal search strategy then takes a very simple form: rank the combinations from the one with the highest expected probability of success ($z = 1$) down to the lowest ($z = 2^N$). Provided combinations have $q_i \geq c$, continue testing each combination down the ranking until a success is achieved and then stop. If no success is achieved when all the combinations with $q_i \geq c$ have been tested, then stop the search. In our drug-search example, the innovator will have a ranked list of potential compounds – or drug designs – with some maximum number that are economic to test. They will then move down the ranking by conducting costly tests, completing the search when either a success is found, or the set of potential economic tests is exhausted.¹³

We assume for simplicity that each combination has a distinct probability of success. To produce the ranking, we define a bijective ranking function, $r(q_i)$, that identifies a rank, z , for each combination in the set based on that combination's expected probability of success,

$$r: q_i \rightarrow z \quad \forall q_i \in \{q_1, \dots, q_{2^N}\} \quad (1)$$

In Fig. 2, we show this bijection for an example with $N = 3$, so that $2^N = 8$, but we note that the actual number of combinations to be ranked could be in the billions. Since r is a bijection, the inverse function exists that maps each element in the set of ranks to a unique probability

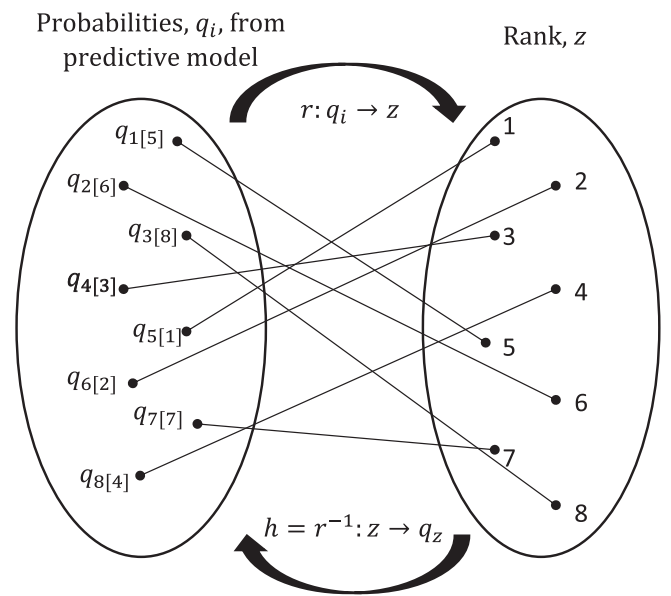


Fig. 2. Bijective function linking probabilities from the predictive model to their rank ($N = 3$; $2^N = 8$). Note: The q_i s on the left are the probabilities of success generated by the prediction model, where the subscript (outside the square bracket) is an arbitrary number indexing a particular design. The designs are then ranked based on probability of success, with the design with the highest probability given a rank of 1 (see the ranking of combinations on the right). The numbers in square brackets on the left represent the re-indexing of the probabilities based on the generated rankings.

of success.¹⁴ Re-indexing the probability of a combination by its generated rank we obtain,

$$h = r^{-1}: z \rightarrow q_z \quad \forall z \in \{1, \dots, 2^N\} \quad (2)$$

This function gives the conditional probability that a combination will yield a success in a test at rank z , where the conditioning refers to

¹⁰ With heterogeneous values and costs, the combinations would be ranked based on their Weitzman reservation prices (see Appendix A.1).

¹¹ We could relax this assumption by assuming that the false positive rate on a test is α , but the market payoff to a successful combination is v and the payoff to a combination that passes the test but ultimately proves unsuccessful in the market is $-\theta v$. We then assume v takes a value such that $(1 - \alpha)v - \alpha\theta v = 1$, or $v = 1/[1 - \alpha(1 + \theta)]$. The expected payoff of a positive test is again equal to 1.

¹² We highlight one important difference from the Weitzman model. Weitzman (1979) assumes the *ex ante* probability distributions for distinct boxes are independent. However, we assume there is a correlation structure across boxes (or combinations in our setting) in a possibly “rugged” search landscape. The ruggedness of the landscape is central to the motivation for using AI to produce predictive models, as the correlation structure may be poorly understood from theory, simple observation, or intuition. However, our assumption of fully exploitive search means that in a given discovery problem the innovator takes the probability distributions over given combinations as given, but is using the available model of the correlation structure – i.e., the predictive model that produces these probability distributions – to prioritize the search.

¹³ It is possible that the entire design space could be tested without a success being found, but we assume that this constraint is not binding.

¹⁴ When probabilities are not distinct, the ranks of two or more combinations can be randomly assigned in the appropriate part of the ranking. For example, if the three combinations that follow rank = 10 all have the same probability of success, they could be randomly assigned to ranks 11, 12 and 13. The mapping from probabilities to ranks is then not a function. However, we can still define a (surjective) conditional success function, where each rank (following any necessary random assignment) maps to a single probability. In this case, we refer to a ranking relation rather than a ranking function, but the subsequent survival analysis based on the conditional success function is unaffected.

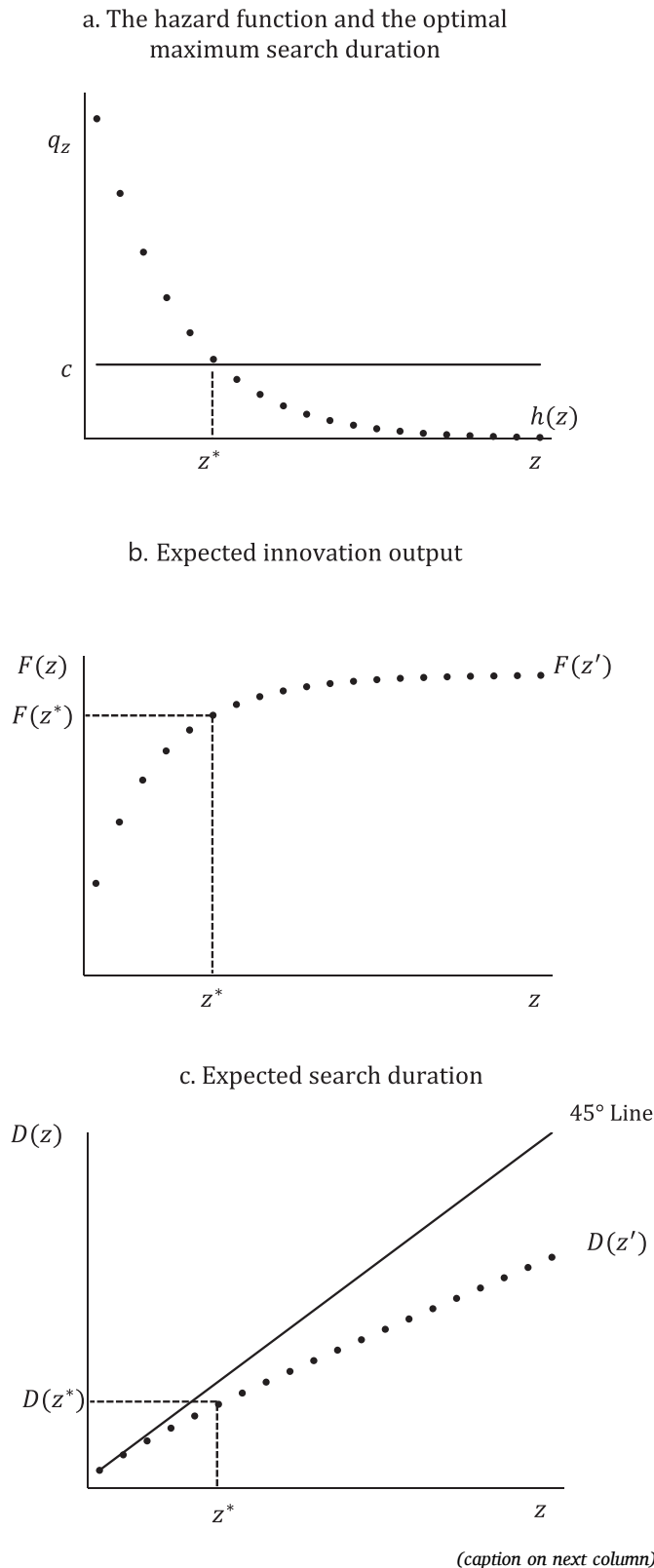


Fig. 3. Key outcomes from an innovation search. Notes: (i) Panel a. shows the probability of success by rank. This can be viewed as a discrete hazard function as it gives the probability of success conditional on reaching that rank in a sequential search. Absent a success, the innovator will continue to test designs provided the probability of success is greater than or equal to the cost of a test. The maximum duration of the search, z^* , can therefore be identified as the highest z such that $q_z \geq c$. (ii) Panel b. shows the cumulative incidence function (or expected probability of innovation) for different durations of the maximum search. This is our measure of expected innovation output. For a given hazard function, a longer maximum search duration is associated with higher expected innovation output. (iii) Panel c. shows how the expected duration of search varies with the maximum duration of search. For a maximum duration of search great than 1, the expected duration will lie below the maximum duration, where we assume there is some positive probability that a success will be found on any given test.

reaching that place in the ranking. We call it the hazard function to underline the analogy with the familiar hazard function in survival analysis.^{15,16} Together with the payoff from a success and the cost of a test, the hazard function completely represents the information available to the innovator.¹⁷

Letting Z indicate the number in the ranking where the first success is discovered, the hazard function produces the probability that a combination $z = Z$ given that the search has “survived” to that place in the ranking,

$$h(z) = \Pr(z = Z | Z > z - 1) \quad (3)$$

We provide an illustration of such a discrete hazard function in Fig. 3a. At each point in the ranking, the function gives the probability that the combination yields a success in testing conditional on reaching that rank. The probabilities are decreasing with the rank because combinations with higher expected probabilities of success have a lower rank number.

It is now straightforward to use $h(z)$ to identify the optimal maximum duration of the sequential search, z^* , as the highest value of z that has an associated expected probability of success that is greater than equal to the cost of searching a combination (i.e., marginal expected value is greater than or equal to marginal cost),

$$h(z^*) = q_{z^*} \geq c \quad (4)$$

We show the optimal maximum duration of the search in Fig. 3a.

2.3. Discrete survival analysis and the expected innovation outcomes

We next define the discrete survival function as giving the probability of surviving to a given place in the ranking without achieving a

¹⁵ The hazard is a rate rather than a probability in continuous-time survival analysis. However, in discrete survival analysis, the hazard has the interpretation of a conditional probability, where the probability is conditional on the rank.

¹⁶ While the simplest case of exploitive search involves a single estimation of the predictive model, the conditional form of the success function allows for continual re-estimation of the predictive model even under exploitive search as the innovator (hypothetically) moves through the ranking, assuming each previous combination tested yields a failure. Recall that the hazard function gives the probability of a success on a particular test given that no successes have been found on previous tests. The remaining combinations can then be re-ranked based on the updated predictive model.

¹⁷ The hazard function is monotonically declining by construction in the case where combinations are ranked simply based on the probabilities of success that are obtained from the prediction model. However, using Weitzman reservation prices, it is still possible to obtain a ranking of combinations for optimal sequential testing. In Appendix A.1, we show that a monotonically declining conditional success function is not necessary for the analysis of innovation outcomes based on this function.

success. The survival probability at z^* is,

$$S(z^*) = \Pr(Z > z^*) = \prod_{z=1}^{z^*} (1 - q_z) \quad (5)$$

In addition to the conditional probability that a given combination at rank z will yield a success, we also define the probability that the *first* success will be achieved at rank z (i.e., $z = Z$). We call this the probability function and use the notation: $f(z) = \Pr(z = Z)$. Evaluating this probability at z^* we have: $f(z^*) = \Pr(z^* = Z) = \prod_{z=1}^{z^*} (1 - q_{z-1})q_{z^*}$. We confirm the familiar definition of the discrete hazard function as the ratio of the probability that the first success is found at z^* divided by the probability of surviving until $z^* - 1$,

$$h(z^*) = \frac{f(z^*)}{S(z^* - 1)} = \frac{\left[\prod_{z=1}^{z^*} (1 - q_{z-1}) \right] q_{z^*}}{\prod_{z=1}^{z^*} (1 - q_{z-1})} = q_{z^*} \quad (6)$$

where we note that $S(0) = 1$. We also note that $f(z^*) = S(z^* - 1)h(z^*)$.

We next define the cumulative hazard function up to a place in the ranking, z^* , as,

$$H(z^*) = \sum_{z=1}^{z^*} q_z \quad (7)$$

Using (5), we exploit a relationship between the log of the survival function and the cumulative hazard function,

$$\ln S(z^*) = \sum_{z=1}^{z^*} \ln(1 - q_z) = -\bar{H}(z^*) \approx -\sum_{z=1}^{z^*} q_z = -H(z^*) \quad (8)$$

That is, the log of the survival function is approximately equal to the negative of the cumulative hazard function. This approximation will be close when the probabilities are low, which is a reasonable assumption when search is over a vast and complex space. The approximation will prove useful below in interpreting the intuitive meaning of an improvement in the AI and we will also refer to $\bar{H}(z)$ as the cumulative hazard function for ease of exposition, although for our formal results we do not assume the approximation holds.¹⁸

Taking the antilog of (8), we obtain an alternative expression for the survival function,

$$S(z^*) = e^{-\bar{H}(z^*)} \quad (9)$$

Given the maximum size of the search, z^* , the *expected innovation output* (or, simply, the probability of innovation) at the outset of that search is given by the cumulative incidence function,¹⁹

$$F(z^*) = \Pr(Z \leq z^*) = 1 - S(z^*) = 1 - e^{-\bar{H}(z^*)} \quad (10)$$

In our drug search example, this gives the probability (as assessed at the beginning of the search) that the overall innovation search will yield a successful drug design. Noting that the cumulative incidence function can also be obtained by summing over the probability function, $f(z)$, we have the alternative way of writing the expected innovation output,

$$F(z^*) = \sum_{z=1}^{z^*} f(z) = \sum_{z=1}^{z^*} S(z-1)h(z) = \sum_{z=1}^{z^*} e^{-\bar{H}(z-1)}q_z \quad (10')$$

In both formulations, expected output is completely determined by the hazard function. We show the expected innovation output function in Fig. 3b.

We turn next to the expected duration of a completed search where the optimal maximum duration of that search is z^* . For simplicity, we assume that a search takes one unit of time, so we can identify the size of a completed search with the time duration of that search. (We relax this assumption in Section 5.) In our drug-search example, we have the expected duration of a search that can end either through finding a successful drug or by exhausting all the designs that are economical to test ($h(z) \geq c$). Together with the cost of a test, the expected duration is sufficient to determine the expected cost of an innovation search.

The concept of restricted mean survival time (RMST) from discrete survival analysis can be used to identify the expected duration. We draw on a classic result in discrete survival analysis, which shows that the RMST is calculated by summing the *lagged* survival functions up to the maximum duration of search,

$$D(z^*) = \text{RMST} = \sum_{z=1}^{z^*} S(z-1) = \sum_{z=1}^{z^*} e^{-\bar{H}(z-1)} \quad (11)$$

where $-\bar{H}(0) = 0$ since $S(0) = 1$. We provide the proof for (11) in Appendix A.2.

Note that the expected duration of a complete search is fully determined by the hazard function. Again, drawing on the hazard function in Fig. 3a, we show in Fig. 3c how the expected duration of the search varies with the maximum duration of the search.

With expressions for expected innovation output the expected duration of the completed search, we can identify the expected profit. Expected profit is simply the expected innovation output less the expected duration of the search multiplied by the cost of a test,

$$\Lambda(z^*) = F(z^*) - cD(z^*) = 1 - e^{-\bar{H}(z^*)} - c \sum_{z=1}^{z^*} e^{-\bar{H}(z-1)} = \sum_{z=1}^{z^*} e^{-\bar{H}(z-1)}(q_z - c) \quad (12)$$

where the last equality uses (10'). We can see that expected profit can be written as an appropriately discounted value of the gaps between expected marginal value and marginal cost up to the maximum duration of the search.

It can be confirmed that the optimal maximum duration of a completed search, z^* , maximizes expected profit by noting that, from the vantage point of the beginning of the search, the expected marginal value of adding an additional combination to the search is $S(z-1)q_z$ and the marginal cost is $S(z-1)c$. Expected profit is maximized by choosing the largest q_z that has an expected marginal value that is greater than or equal to marginal cost,

$$S(z^* - 1)q_{z^*} \geq S(z^* - 1)c \quad (4')$$

which yields the same optimal maximum duration as (4).

To recap the mechanics of the model in terms of the hazard function notation, we note that the information available to the innovator from the predictive model is captured in the form of $h(z)$ and its associated $\bar{H}(z)$. $h(z)$, together with the exogenous cost of a test, is used to determine the maximum duration of a completed search; with this maximum duration identified, the key innovation outcomes can be calculated from knowledge of $\bar{H}(z)$. The centrality of the discrete hazard function to the innovation outcomes also provides a convenient way to specify improvements in the predictive model based on increases in the cumulative

¹⁸ If the domain of the success function is treated as continuous rather than discrete, then the log of the survival function will be identically equal to the negative of the cumulative success. Although vast, it is more natural to treat the search space over all potential combinations as discrete rather than continuous, but the discrete assumption does come with a small cost in analytic tractability.

¹⁹ We do not refer to this as a cumulative distribution function as it is possible that no success is found even after a fully exhaustive search. We therefore have $\sum_{z=1}^{z^*} f(z) \leq 1$, where z^* is the total size of the search space.

hazard function over some relevant range of ranks.²⁰

As noted above, we have assumed the innovator has access to an unbiased prediction model that is estimated using an IID-drawn sample of successes and failures from the underlying combinatorial space of potential designs. In real world applications, the innovator is likely to face significant challenges in estimating a prediction model that performs well out of sample, with consequent implications in the reliability of the estimated hazard function for guiding the sequential search. Actual applications will involve the use of judgement to balance the inevitable bias-variance trade-off in estimating the prediction model, although we assume the innovator will always choose the prediction model based on its performance in a held out sub-sample of the training data. However, we underline that a significant potential limitation of our analysis is that it presumes that an unbiased prediction model is available.

The idea of modelling the innovation as a process of search is not new. One particularly tractable formulation is to model innovation search as a Poisson process, where the per-period probability of success conditional on not finding a success in previous periods is fixed (i.e., the success hazard rate is constant) but the timing of the first success is random. Dasgupta and Stiglitz (1980), for example, model R&D as a Poisson process where the per-period probability of success depends on R&D expenditure.²¹

Hazard models also have a distinguished history in the economic modelling of innovation. However, these models typically take the form of a black box where a given level of investment in R&D leads to a constant hazard. Discovery is then certain over an infinite horizon (with the cumulative hazard going to infinity) and the expected time-to-discovery is equal to one divided by the hazard rate. Our model allows for a richer treatment of the hazard function and a completed search need not result in a discovery. This results from our modelling of optimal sequential search, where search requires costly tests, but a predictive model is used to prioritize those tests. We can then capture an improvement in the predictive model in terms of its implied effects on the hazard function.

Another fruitful approach in the literature is to model innovation search as draws from a known distribution, and to derive the implied extreme value (best draw) distribution. Kortum (1997) shows how an appropriate rate of increase in sampling from an (thick-tailed) Pareto distribution is consistent with exponential growth in the technological frontier representing the best available technology. In a recent paper that also treats innovation search as draws from a distribution, Jones (2021) identifies conditions under which combinatorial growth in the number of possibilities leads to exponential economic growth. In contrast to Kortum (1997), the assumption of combinatorial growth in the number of possibilities is sufficient to generate exponential growth even for thin-tailed underlying distributions. While ingenious, this formulation depends on exhaustive evaluation of the potentially vast set of possibilities. In our approach, the emphasis is instead on how the search is prioritized using a predictive model, which we argue provides a better way to conceptualize observed uses of AI as a tool for scientific discovery and innovation where testing is costly.

In the next section, we consider when a well-defined improvement in the prediction model will lead to “better, faster, cheaper” innovation.

²⁰ In Appendix A.1, we relax the assumptions of homogeneous success payoffs and costs of testing across design combinations, and show that the key result established above holds true: given the payoffs and costs, the conditional success function (and associated cumulative success function) provide sufficient information to determine the key innovation outcomes.

²¹ Aghion and Howitt (1992) model the flow of new ideas as a Poisson process in developing a quality-ladder model of endogenous growth based on creative destruction.

3. Improvements in prediction and “Better, Faster, Cheaper” innovation

We now consider the effects of an improvement in the predictive model, initially assuming that the maximum number of hypotheses tested (combinations sent for testing) by the innovator stays at z^* . To be concrete, we assume that this improvement comes about because of a newly available AI-based prediction model.

The improved predictive model generates more discriminating predictions, and is therefore assumed to increase the probability of success of the well-ranked combinations. Notwithstanding that the rank order of the combinations may change, we intuitively expect that an improvement in the prediction model causes the hazard function to swivel in a clockwise direction, thereby causing the cumulative hazard function to rise over some range of the best-ranked combinations.²²

Under what conditions will the search yield higher expected innovation output and a lower expected search duration for a given maximum duration of search? Together these outcomes would also imply unambiguously higher expected profit.

The following conditions on $\bar{H}(z)$ are *sufficient* to yield an innovation process that is better (higher expected output and profit), faster and cheaper (lower expected search duration and consequent cost):

$$\bar{H}_1(z) > \bar{H}_0(z) \forall z = 1, \dots, z^* \quad (13)$$

where 0 and 1 index the prediction model before and after the improvement in the prediction model respectively. We can intuitively think of these conditions as saying that the improvement in the AI leads to a higher cumulative hazard function at each rank up to z^* .²³

We can use (10) and (11) to determine the effect of (13) on expected innovation output, the expected duration search, and (putting these together) on expected profit. Given the assumed improvement in the prediction model: $\Delta F(z^*) > 0$, since $\bar{H}_1(z^*) > \bar{H}_0(z^*)$; $\Delta D(z^*) < 0$, since $\bar{H}_1(z) > \bar{H}_0(z)$, $\forall z = 1, \dots, z^* - 1$; and therefore, $\Delta \Lambda(z^*) > 0$.

The *necessary and sufficient* conditions for expected innovation output to increase and expected duration to fall are: $\bar{H}_1(z^*) > \bar{H}_0(z^*)$ and $\sum_{z=1}^{z^*} e^{-\bar{H}_1(z-1)} < \sum_{z=1}^{z^*} e^{-\bar{H}_0(z-1)}$, where the latter condition requires just that the value of the cumulative lagged survival probability is lower at z^* , rather than requiring that the lagged survival probability is lower at each rank up to and including z^* .

The foregoing analysis assumes that, following the improvement in the prediction model, the maximum number of tests that could be completed before the search is abandoned is unchanged at z^* . However, the improved prediction model will lead the innovator to re-optimize in terms of choice of the *maximum* number of combinations to send for testing, with the maximum number of tests potentially either rising or falling depending on the change to the shape of $h(z)$. We denote the new

²² To explain why we refer to an improvement in the prediction model as leading to a clockwise swivel of the hazard function curve, it is useful to consider a hypothetical case where there is a given number of successful combinations in the design space. With perfect prediction, the hazard function would take the shape of a unit step function, with a probability of success equal to 1 for all the actual successes in the design space and a value equal to 0 for all the actual failures. At the other extreme, where the prediction model has zero predictive power, the hazard function would be a horizontal line at value equal to the probability of success of a random draw from design space. Starting from a given prediction model, we can therefore usefully think of an improvement in the prediction model (or move towards a more “discriminating” prediction model) as a move in the direction of the unit step function – i.e., it will move in way that represents a clockwise swivel of the hazard function curve, raising the probability of success of highly ranked combinations and lowering the probability of success of poorly ranked combinations.

²³ These are weaker conditions than requiring this probability to be higher at each of these ranks, and allows in principle for the new conditional success function to cross the old success function before z^* .

optimal maximum search as z^{**} .

The new optimal maximum search is given by the intersection of the hazard function (post prediction-model improvement) with the test cost line. It is possible that the new optimal maximum number of tests could either rise or fall. An improved prediction model could increase the number of combinations that it is economic to test ($z^{**} > z^*$). For example, for an innovator searching for a small molecule drug to bind with a target protein, a greater range of possible designs may become economic to test with improved knowledge of the design space. But to the extent that the improved discrimination that is possible with a better predictive model – raising the probability of success of the best ranked combinations, but lowering that probability for most others – the swiveling of the hazard function might be such that the optimal maximum number of tests actually falls ($z^{**} < z^*$). For example, the innovator seeking the new small molecule drug might be better able to effectively rule out a larger fraction of the design space.

For $z^{**} \neq z^*$, the necessary and sufficient conditions for expected innovation output to increase and expected duration to fall are then, respectively: $\bar{H}_1(z^{**}) > \bar{H}_0(z^*)$ and $\sum_{z=1}^{z^{**}} e^{-\bar{H}_1(z-1)} < \sum_{z=1}^{z^*} e^{-\bar{H}_0(z-1)}$.

To help fix intuition, in Fig. 4 we illustrate two cases involving an improvement in the prediction model. In both cases, we assume that the cumulative hazard function is larger following the improvement when measured at z^* .

In Fig. 4a we show the case where the new hazard function crosses the cost line at $z^{**} > z^*$. Moving to this new profit maximizing maximum duration of search will result in a further increase in expected innovation output and profit, but the effect on the expected duration of search will only be negative (i.e., faster innovation) if $\sum_{z=1}^{z^{**}} e^{-\bar{H}_1(z-1)} < \sum_{z=1}^{z^*} e^{-\bar{H}_0(z-1)}$.

In Fig. 4b, we show the case where $z^{**} < z^*$. In this case, expected profit will again increase and expected duration will decrease. But the effect on expected innovation output will now only be positive if $\bar{H}_1(z^{**}) > \bar{H}_0(z^*)$.

In summary, the linkage of the prediction model and its associated ranking function with survival analysis allows us to build a tractable model of search over a combinatorial search space. In particular, we can represent the information gain that results from an improvement in the prediction model as a change in the position of the hazard function and the associated change in the cumulative hazard function. Intuitively, we can think of an improvement in the AI as yielding an increase in the cumulative hazard function over the relevant range of ranks. Assuming optimal sequential search, we can identify general conditions on the cumulative hazard function under which the availability of AI leads to “better, faster, cheaper” innovation.

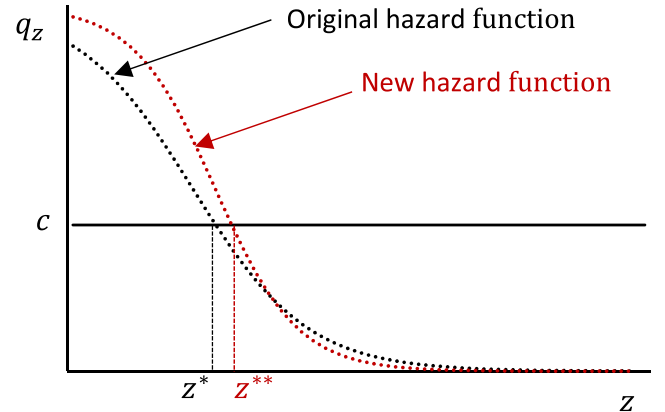
4. Extensions to baseline model

The baseline model makes a number of simplifications in order to highlight the essential features of predictive-model-aided search over a combinatorial search space: no heterogeneity in success outcomes; no competition between innovators in seeking to discover a successful design; and no exploration to enhance the predictive model. These assumptions obviously abstract away from important features of real world innovation processes. Given space limitations, we simply note how the basic model can be extended to incorporate these additional elements.

4.1. Heterogeneity in success outcomes

Although it significantly simplifies the analysis, a restrictive feature of the baseline model is homogeneity of payoffs conditional on a combination being a success. One limiting feature of this assumption is that it excludes the relative “riskiness” of a combination being relevant for its ranking, which is at odds with an intuition that innovators may

a. Improvement in the prediction model where the optimal maximum duration of search increases



b. Improvement in the predictive model where the optimal maximum duration of search decreases

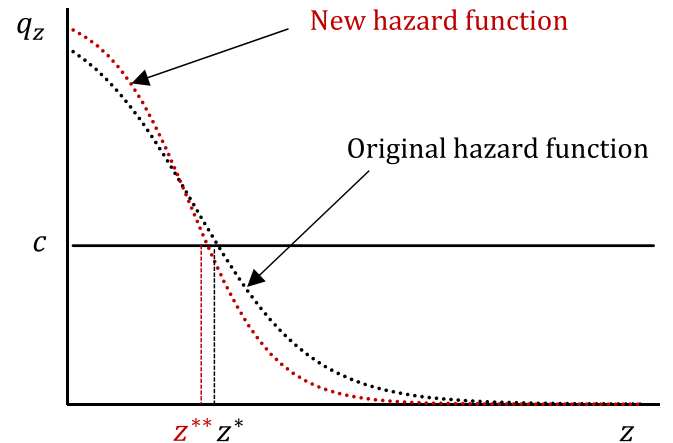


Fig. 4. Impact of an improvement in the prediction model (assumes cumulative conditional success probability increases at the original profit maximizing maximum search duration). Notes: An improvement in the prediction model is represented as a clockwise swivel in the hazard function, reflecting an increase in the probabilities of success of the highly ranked combinations over some range due to the more discriminating prediction model. Panel a. shows the case where the maximum duration of search increases; Panel b. shows the case where the maximum duration of search decreases. In both cases, given the optimal maximum duration of search, the pre- and post-improvement innovation outcomes can be obtained from knowledge of the cumulative hazard function over the relevant ranges of ranks as set out in the text.

prioritize high payoff but low probability potential combinations in the innovation search. However, even with Bernoulli boxes, the Weitzman model allows for heterogeneity in payoffs (and costs) as well probability of success to be relevant for the calculation of “reservation prices” underlying the optimal order of search.

In Appendix A.1, we relax the assumption of common payoffs conditional on success and show that the main results of the basic model still hold. The more general model shows that high payoff/low probability potential combinations will be prioritized. We show that for two

potential combinations (or boxes) with the same expected value, $q_1 v_1 = q_2 v_2$, but where $v_1 > v_2$ and $q_1 < q_2$, Box 1 will have a higher reservation price than Box 2 and consequently will be searched first. As there is an intuitive sense that Box 1 is riskier than Box 2 – a higher payoff conditional on success but a lower conditional probability of success – this shows how an innovator will tend to prioritize riskier combinations in the search.

Moving beyond Bernoulli boxes, the general form of the Weitzman model allows the probability distribution of payoffs to be specified on the continuous domain $(-\infty, \infty)$. It is noteworthy that Weitzman emphasizes the tendency for “riskier” options to be searched first in an optimal sequential search as one of most important implications of the model:

Other things being equal, it is optimal to sample first from distributions that are more spread out or riskier in the hopes of striking it rich early and ending the search. This is a major result of the present paper. Low-probability high-payoff situations should be prime candidates for early investigation even though they may have a smaller chance of ending up as the source ultimately yielding the maximum reward when the search ends. (Weitzman, 1979, p. 647.)

4.2. Competition between innovators

A second restrictive feature of the baseline model is that it ignores potential competition between innovators to discover a success – say an effective drug to treat Alzheimer’s disease. Even without explicit strategic interaction between competing innovators (as in a classic “patent race”), innovators may choose their discovery problem taking into account the competition. Given our use of survival analysis in the model, this suggests the “competing risks” paradigm from epidemiology as a natural way to extend the model to allow for competition.

Suppose, for example, there are K competitors ($k = 1, \dots, M$) each using a distinct hazard function. All competitors start their searches at the same time. We further assume that owing to the size of the search space, the prioritized combinations are not overlapping in the relevant range of the functions. The innovation search will end if any competitor achieves a success. The cumulative incidence function for our focal innovator, j , then takes the form:

$$(z) = \sum_{z=1}^{z^*} \left(\sum_{k=1}^K S_k(z-1) \right) h_j(z) \quad (14)$$

where $S_k(\cdot)$ is the survival function for competitor k and $h_j(z) = q_{j,z}$ is the probability of success for the focal innovator conditional on reaching rank z .

In choosing an innovation problem, the innovator may face a choice between problems for which a good prediction model exists but there are a large number of competitors, and problems with a less accurate prediction model but also fewer competitors in the race for discovery.

4.3. Allowing for an exploration-exploitation tradeoff

The baseline model assumes the innovator takes the predictive model as given and simply exploits it to maximize the expected profit from innovation. However, an innovator may also choose tests in order to generate data to improve the predictive model rather than simply exploiting the predictive model it has. Such an exploration motive is strengthened if the innovator is involved in multiple projects and there is a possibility of transfer learning between projects, whereby the data generated in the focal project can improve the predictive models used in their other projects.

We consider two possible directions to extend the model that allow for exploration as well exploitation. First, we assume that there is an initial exploration stage focused on improving the predictive model followed by a standard exploitation stage. To keep the stages distinct, we

assume that exploration takes place on an adjacent problem and transfer learning can be used to improve the predictive model for the focal problem.

The vast literature on *active learning* points to a fruitful way of modelling the exploration stage. Instead of taking the data available to train the predictive model as given, the innovator can request a “label” from an “oracle” – i.e., the test of a target combination in our setting. Various criteria for selecting the tests have been suggested in the active learning literature. As examples, under *uncertainty sampling*, tests would be requested based on the degree of uncertainty surrounding particularly labels; under *diversity sampling*, tests would be requested for regions of the search space where data is particularly sparse. Recognizing these tests are costly, we assume there is an optimal set of tests that would be performed in this exploratory stage, which we can think of as a fixed cost for the innovation process.

While the two-stage process neatly separates the exploration and exploitation stages, the nature of sequential search suggests the importance of ongoing exploration, so that the exploration and exploitation take place in parallel. The paradigmatic example of an exploration-exploitation tradeoff is the multi-armed bandit problem, but the standard formulation is not suitable in our setting given the assumed correlation structure in the search space and the fact that search ceases once a success is found.

However, the broad approach of *contextual bandits* suggests a possible framework for modelling the tradeoff. The relevant contexts are the features of the potential combinations (e.g., molecular compounds in the small drug discovery example) and the predictive model maps a given context to its probability of success. For any given ordering of tests, we can still identify the relevant hazard function, but it is now possible that lower probability of success combinations will be ranked better than higher probability combinations. Intuitively, a lower probability combination may receive a better ranking in part because of the informativeness of the combination for the predictive model, allowing for an “improved” hazard function over later parts of the ranking.

In effect, we can view our innovator as choosing over different hazard functions, where each function reflects a different exploration-exploitation tradeoff. More formally, we can define a functional relationship: $\Lambda = \Lambda(h(z))$, which maps the expected profit from the search to the particular hazard function, where $h(z) \in \Theta$, where Θ is the set of feasible hazard functions given the exploration-exploitation tradeoff.²⁴ The optimal hazard function is then:

$$h(z)^* = \underset{h(z) \in \Theta}{\operatorname{argmax}} [\Lambda(h(z))] \quad (15)$$

And the optimal expected profit is:

$$\Lambda^* = \Lambda(h(z)^*) \quad (16)$$

While this optimal hazard function exists in principle given our assumptions, in reality the contextual bandit literature suggests various heuristic methods will have to be deployed to achieve reasonable performance on the exploration-exploitation tradeoff.

5. Autonomous discovery systems

To the extent that AI provides prediction models that prioritize

²⁴ A restricted example would be where the innovator first engages in exploratory search by randomly selecting a given number of designs from the design space to undertake tests to add to the data available to estimate the predictive model. In advance of the exploration, the innovator will not know the precise ranking of the combinations they will face when they enter the exploitation stage. But we assume that they can determine the position of hazard function in the exploitation range. We can therefore think of the innovator as choosing between different hazard functions given the set of exploration strategies available.

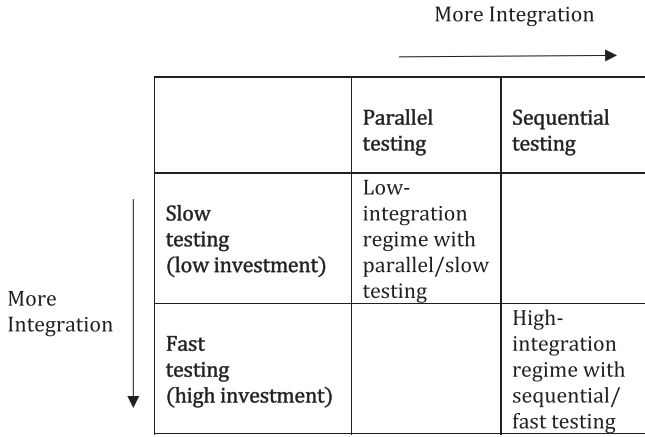


Fig. 5. Complementary dimensions of Integration in the Innovation Process. Notes: The diagram captures an innovation system with two complementary dimensions: whether testing is done in parallel or sequentially; and whether testing is slow or fast. It is assumed that parallel/slow and sequential/fast represent two coherent systems. An investment option is available that affects the speed of testing. It is assumed that the sequential/fast system is superior given optimal investment (see text). However, we allow for the possibility that the innovator will remain stuck in the inferior system given difficulties of coordinating the reorganization required to move from a parallel to a sequential testing regime together with the investment in speeding up the testing process.

search over various combinatorial search spaces – genes, proteins, molecules, etc. – it can be viewed as a meta-GPT. A major theme of the theoretical and historical literatures on GPTs is that significant reorganizations are often required before the new GPT has its full impact. However, these reorganizations can be held up by the need to coordinate different elements of the system change.²⁵

In this section, we extend and apply the baseline model to the question of system-level reorganization. We assume there are two dimensions to the prediction-model-aided innovation system: whether testing is done in parallel or sequentially; and whether testing is slow or fast. We show that parallel/slow and sequential/fast can both form coherent systems (Fig. 5). We also allow for the sequential/fast system to ultimately be more profitable, although the transition to the superior system can be delayed due to coordination challenges associated with system redesign.

We first extend the model to allow for parallel search and compare the innovation outcomes with our baseline sequential search model. We then introduce a test-time variable that is absent in the baseline model (where all tests were assumed to take one period). Finally, we consider the extreme case of a fully autonomous innovation system.

5.1. Parallel versus sequential testing

We have assumed to this point that the testing process is sequential. However, where the testing process is lengthy (e.g., the multiple stages involved in testing for the safety and efficacy of a new drug), the innovator may be forced to test different potential combinations in parallel rather than in sequence. We model this as the innovator choosing a *portfolio* of combinations at the outset of the search that proceed in parallel through the testing process.

²⁵ This form of process re-organization has parallels in the evolving responses to other important GPTs. In David's (1990) classic treatment of the introduction of electricity, substantial productivity benefits were seen only once factories were re-organized – notably from a vertical design with a central power source to a horizontal design with distributed power sources – which can be viewed as a movement towards greater integration with the new GPT. See also Agrawal et al. (2022).

The problem facing the innovator is to choose the size and composition of this portfolio. In hypothetically adding a combination to the portfolio, the innovator must consider the probability that a success will be achieved by one of the already included members. With m existing members, the marginal expected value of adding a member is then: $f(m+1) = (1 - q_1)(1 - q_2) \dots (1 - q_m)q_{m+1}$. This is the probability function evaluated at $m+1$ – i.e., the probability of achieving the *first* success on combination $m+1$. To obtain the optimal portfolio size (and composition), we replace the function with the probability function in (4). The optimal size of the portfolio, z^+ , is then given by the highest ranked combination such that,

$$f(z^+) \geq c \quad (17)$$

To compare parallel and sequential testing, it is useful to rewrite this optimality condition as,

$$f(z^+) = S(z^+ - 1)q_{z^+} \geq c \quad (18)$$

Comparing (18) to (4'), it is clear that $z^+ < z^*$ at the respective profit-maximizing points; i.e., the size of the optimal portfolio under parallel testing is less than the size of the maximum search under sequential testing. The reason is that although the respective expected marginal value schedules are the same, the expected marginal cost schedule is lower under sequential testing given the less than certain probability of reaching any given rank ($1 < z \leq z^*$) under sequential testing. Moreover, if we force the size of the maximum search under sequential testing to be z^+ , expected profits and duration will be lower as the search will stop before reaching z^+ if a success is found, while the expected innovation output will be the same. By allowing the maximum length of the sequential testing to be chosen optimally, expected profits from sequential testing would rise further, as would expected innovation output, although it is possible that the expected duration of the sequential testing would be longer than the optimal parallel testing. The expected profit under parallel testing is,

$$\begin{aligned} \Lambda(z^+) &= \sum_{z=1}^{z^+} S(z-1)q_z - \sum_{z=1}^{z^+} c \\ &= \sum_{z=1}^{z^+} e^{-\bar{H}(z-1)} q_z - cz^+ < \Lambda(z^*) \end{aligned} \quad (23)$$

We show profit as a function of different search lengths (representing the size of the portfolio under parallel testing and the maximum duration of the search under sequential testing) in Fig. 6. The bottom line of this comparison is that, with no penalty for longer searches, an increase in integration as reflected in a move from parallel to sequential testing will result in higher expected profits and innovation output.

5.2. Slow versus fast testing

Of course, the reason for the greater profitability of sequential testing is that there is no penalty for the extra time taken to engage in sequential compared to parallel testing. But if testing takes time, parallel testing has the inherent advantage that the entire portfolio can be tested in the time it takes one test under sequential testing.

To provide a fair comparison between the two forms of testing, we therefore need to introduce a time penalty, which we do simply by introducing time discounting. Starting at time $t = 0$, future benefits and costs are discounted by the discount factor δ^t . Furthermore, we now assume that a test takes λ units of time. We can therefore rewrite the discount factor in units of the test as $\delta^{\lambda z}$.

Under sequential testing, for a given rank in the sequence that a test results in success, the higher the value of λ the more the resulting expected profit will be discounted. As the entire portfolio can be tested under parallel testing in the time it takes to conduct one test, the relevant discount factor is a constant δ^λ across the entire portfolio.

Allowing for discounting, the optimal discounted expected profits

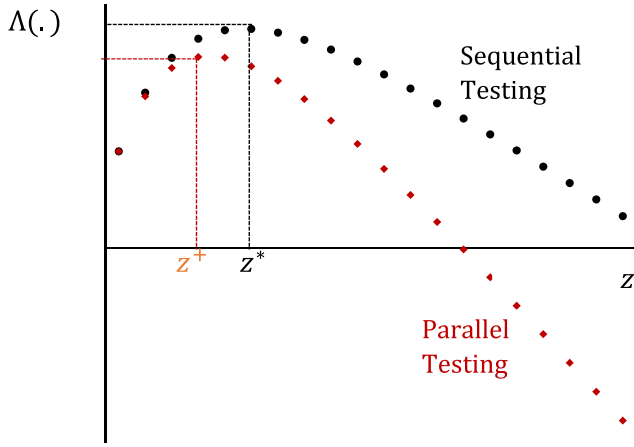


Fig. 6. Expected profit under parallel and sequential testing. Note: The figure shows how the profitability of sequential and parallel testing varies with the maximum duration of testing (sequential testing) and the size of the testing portfolio (parallel testing). The optimal maximum duration is z^* (sequential testing) and the optimal portfolio size is z^+ (parallel testing). With no time penalty, maximum profit is higher under sequential testing than under parallel testing.

under sequential and parallel testing are, respectively:

$$\Lambda(z^*) = \sum_{z=1}^{z^*} \delta^{\lambda z} S(z-1)(q_z - c) \quad (24)$$

$$\Lambda(z^+) = \sum_{z=1}^{z^+} \delta^{\lambda z} (S(z-1)q_z - c) \quad (25)$$

It is useful to write the gap between profits under sequential and parallel testing as,

$$\Lambda(z^*) - \Lambda(z^+) = \delta^{\lambda} \left[\left(\sum_{z=1}^{z^*} \delta^{\lambda(z-1)} S(z-1)(q_z - c) \right) - \sum_{z=1}^{z^+} (S(z-1)q_z - c) \right] \quad (26)$$

As noted, when $\lambda = 0$ this gap is positive (i.e., sequential testing is more profitable than parallel testing). However, as λ increases the gap will decrease as the term in square brackets becomes smaller and will eventually become zero.

Again, using the conditional success function from Fig. 3b, in Fig. 7 we compare the optimal expected profits under sequential and parallel testing for various values of λ . When $\lambda = 0$ we obtain the same optimal profits as with no discounting since testing is instantaneous and time is irrelevant. However, the figure shows the increasing penalty of relatively time-consuming sequential testing as λ increases, and the implied optimal regimes of sequential and parallel testing.

As λ is a continuous variable, there is a value of $\lambda = \tilde{\lambda} > 0$ that makes the innovator indifferent between sequential and parallel testing. We can therefore identify two regimes: (i) a sequential testing regime, $0 \leq \lambda < \tilde{\lambda}$; and (ii) a parallel testing regime, $\lambda \geq \tilde{\lambda}$, where we assume (arbitrarily) that the innovator uses parallel testing in the case of indifference. Moreover, the size of the sequential testing regime will shrink with the size of the discount rate, δ , as the innovator becomes more sensitive to the timing that profit is realized.

Our next step is to endogenize λ and thus the innovator's choice of the testing regime. For any given starting value of the testing time, λ_0 , we assume the innovator can make investments $I = I(\lambda_0 - \lambda)$ to reduce the length of testing time. The cost of the investment, I , is assumed to be increasing at an increasing rate in the size of the reduction in the testing time; i.e., the marginal cost of reducing the testing time is increasing with the size of the reduction. Under either regime, the marginal benefit

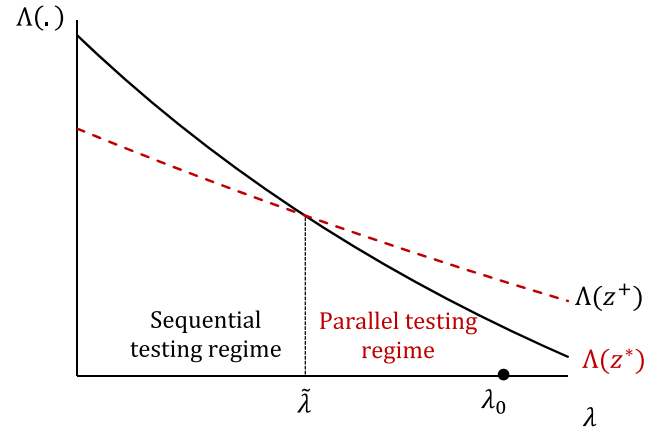


Fig. 7. Sequential and parallel testing regimes. Note: Based on the hazard function in Fig. 3b this figure shows expected profit under optimal sequential and optimal parallel testing as a function of the time cost of conducting a test. There is a cut-off time cost where the expected profitability is equal under both regimes. Sequential testing yields higher expected profit for time costs below this cut-off; parallel testing yields higher expected profit for time costs at or above this cut-off. We can therefore identify optimal sequential and parallel testing regimes based on the actual time cost the innovator faces.

is the marginal increase in *optimal* regime-specific profit: $-\partial\Lambda(z^*)/\partial\lambda$ under sequential testing; and $-\partial\Lambda(z^+)/\partial\lambda$ under parallel testing. In Fig. 8, we show the assumed marginal cost curve and the two marginal benefit curves for our illustrative example based again on the empirical conditional success function shown in Fig. 3a.

Under each regime, we can identify the optimal testing time as,

$$\text{Sequential Testing : } \lambda^* = \underset{\lambda < \tilde{\lambda}}{\operatorname{argmax}} [\Lambda(z^*, \lambda) - I(\lambda_0 - \lambda)] \quad (27)$$

$$\text{Parallel Testing : } \lambda^+ = \underset{\lambda > \tilde{\lambda}}{\operatorname{argmax}} [\Lambda(z^+, \lambda) - I(\lambda_0 - \lambda)] \quad (28)$$

The optimal *regime* is then given as,

$$\lambda^* = \underset{\lambda^* \leq \tilde{\lambda}}{\operatorname{argmax}} [(\Lambda(z^*, \lambda^*) - I(\lambda_0 - \lambda^*)), (\Lambda(z^+, \lambda^+) - I(\lambda_0 - \lambda^+))] \quad (29)$$

In terms of Fig. 5, recognizing the likely complementarities between sequential testing and speedy testing, the innovator is most likely to be

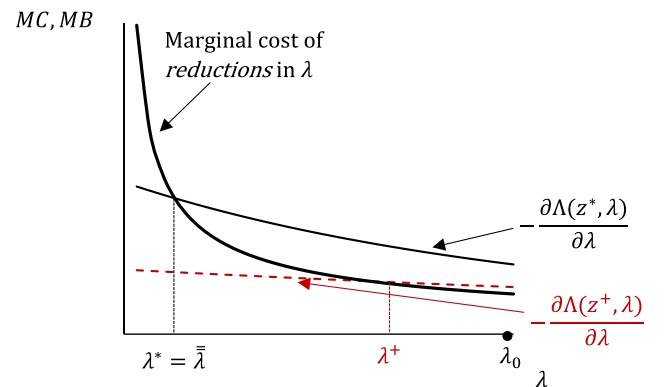


Fig. 8. Endogenous determination of testing regime. Note: The figure shows the expected marginal benefit and marginal cost of small reductions in the test time relative to some starting test cost, λ_0 . The optimal test cost is λ^* under sequential testing and λ^+ under parallel testing. Given optimal investment under each regime, it is assumed in the figure that expected profit is higher under the sequential testing regime. Therefore, the overall optimal test time, λ , is λ^* .

in the upper left-hand quadrant or the lower right-hand quadrant. We think of an improvement in the prediction model as potentially shifting the innovator (after some reorganization involving costly investments) from the parallel/slow to the sequential/fast regime. Given the needed reorganization, time could elapse between the availability of a new GPT for prediction and visible changes in the testing process and outcomes.

5.3. An autonomous discovery system

We chose these two dimensions of integration for their plausibility but also because they can be demonstrated with straightforward extensions of our model. Other possibilities arise as we move towards greater integration via rapid sequential testing. As discussed in [Section 4](#), more integrated search processes would also give rise to possibilities for more effective *exploration* of the search space, which may have benefits beyond the immediate design task. Such exploratory search would be aimed at improving the predictive model, and the benefits would also include the information gains, some of which may spillover to other design problems and other innovators.

We, therefore, conceive an autonomous (or “self-driving”) testing process as having the following complementary characteristics: (i) sequential testing; (ii) rapid testing; (iii) a predictive model that effectively prioritizes tests over a wide range of combinations; and (iv) feedback from the tests to the predictive model allowing for exploratory search.

Materials science provides a possible use-case for such autonomous discovery systems. As with drug discovery, the space of potential molecules is vast. Computational methods have long been used to aid the discovery process, including the use of computational chemistry to virtually screen for the properties of molecules, including methods such as quantum chemistry and molecular mechanics. However, the computational costs of such simulation methods can be prohibitive, leading to interest in statistical approaches such as machine learning to prioritize molecules for simulation- or experiment-based characterization (see, e.g., [Pyzer-Knapp et al., 2015](#)). For example, input data on molecular descriptors and output data on molecular properties could be used to develop a machine-learning-based predictive model of a large chemical space that would otherwise be prohibitively costly through computational and experimental methods.

Lamenting the slow speed and high cost of the development and deployment of advanced materials using the traditional approach – new materials typically “reach the market after 10–20 years of basic and applied research” – [Tabor et al. \(2018, p.5\)](#) outline what they see as required for an autonomous (closed-loop) innovation process:

To fully exploit the advances in autonomous robotics, machine learning, high-throughput virtual screening, combinatorial methods and in situ or in operando characterization, we must close the loop in the research process. This means that humans must partner with autonomous research robots to design experimental campaigns and that the research robots perform experiments, analyze the results, update our understanding and then use AI and machine learning to design new experiments optimized to the research goals, thus completing one experimental loop.

Such autonomous discovery systems are already under active exploration ([Aspuru-Guzik and Persson, 2018](#)). These systems are feasible where it is possible to have rapid feedback in terms of outcome data, where those data can be used to improve the prediction model as part of, say, an active learning strategy. For example, machine-learning-based predictions determine which candidates will be tested next using robotic high throughput screening (HTP) methods.

There is an obvious potential for a move to a capital-intensive, self-driving discovery process to alter the balance between humans, algorithms, and robots in the R&D process. The most obvious threat in terms of demand for task-specific skills is to those involved in testing that is increasingly automated. However, chemists involved in selecting the

candidates for testing (hypothesis generation) would also be at risk, as AI-based predictions replace human predictions in the prioritization process. On the other hand, new demands would likely arise in building task-specific AIs for prediction. To guide the overall design process, there may also be increased demand for individuals who can integrate AI and domain-specific skills, such as AI specialists who have an understanding of chemistry or chemists with an understanding of AI.

6. Discussion

Our model suggests that one testing regime, fast and sequential, might be superior, from a social welfare perspective, compared to the other regime, slow and parallel. However, it's possible to get stuck in the slow-parallel equilibrium. The advance of AIs might so dramatically increase the quality of hypothesis generation that this alone might drive the shift from slow-parallel to fast-sequential. Or, this might not. A social planner might be required to intervene and move society to the more socially optimal equilibrium.

Higher fidelity hypotheses increase the return on investment in testing capacity, so as AIs improve, they will increase the incentive to invest in testing. An increase in testing capacity will, in turn, increase the incentive to further invest in AIs for improved hypothesis generation. AIs for hypothesis generation and technologies for hypothesis testing are complements. So, we could enter a renaissance in scientific discovery if the flywheel begins to spin as investments in AI drive investments in testing which drive further investments in AI.

However, the flywheel may not begin to spin, and thus society may not benefit from a flourishing of new scientific discoveries, without policy intervention. How will these complements co-evolve? What is the role for policy? Although the path is uncertain, we can gain policy insights from observing the evolutions of past GPTs as seen through the lens of the model. As introduced by [Bresnahan and Trajtenberg \(1995, p. 84\)](#) GPTs are:

... characterized by the potential for pervasive use in a wide range of sectors and by their technological dynamism. As a GPT evolves and advances it spreads throughout the economy, bringing about and fostering generalized productivity gains... This phenomenon involves what we call “innovational complementarities” (IC), that is, the productivity of R&D in a downstream sector increases as a consequence of innovation in the GPT technology. These complementarities magnify the effects of innovation in the GPT, and help propagate them through the economy.

The special relevance of GPTs that provide a new technology for innovation was recognized early on by Nathan Rosenberg ([Rosenberg, 1998](#)). Writing in the context of breakthroughs in the discipline of chemical engineering (and in particular, the invention of “unit operations”), [Rosenberg \(1998, p. 165\)](#) writes:

This perspective of chemical engineering as a general purpose technology may be compared with GPTs that have been identified with a specific form of hardware: steam engines, machine tools, dynamos, computers, and so on. This suggests that the concept of a GPT should not be confined to hardware. Indeed, a discipline that provides the concepts and methodologies to generate new or improved technologies over a wide range of downstream activity may be thought of as an even purer, or higher order, of GPT.

The theoretical and empirical literatures on GPTs have shown how various market failures can condition the speed and long-run impacts of the technology. First, past GPTs have shown how full impacts are seen only after substantial process reorganizations – such as reorganizing factory floors to take advantage of electricity. Such reorganizations may raise substantial coordination challenges for managers and policymakers. Our treatment of an autonomous discovery system provides an example of how rapid exploratory testing may be required to take advantage of the predictive possibilities of modern AI systems.

Second, the evolution of past GPTs has highlighted the importance of horizontal spillovers. Drawing on our model, two types of spillovers stand out as being particularly important: demonstration effects and data.

Demonstration effects refer to how successes in one design problem can serve as analogies for how to go about the process of discovering designs in other domains. This could happen within sectors – say, the finding of a small molecule drug that successfully binds with a target malfunctioning protein could demonstrate how to find a drug to bind with a protein related to a different disease; and across sectors – say demonstrations of successful use in medical biotechnology could help solve design problems in agricultural biology. In the context of AI applications, such demonstration effects may be more important than in other areas of science given the emphasis on finding something that “works” notwithstanding the often paucity of understanding of “why it works.”

Recognizing the importance of data to modern statistics-based AI, a potentially even more consequential source of spillovers is data.²⁶ In our model, testing generates data on successes and failures, and these data are used to improve the underlying predictive models. Even putting aside for the moment concern about the monopolization of data, scientists may have difficulty accessing data that exists in principle as a non-rival good. Data on what has *not* worked – failures in our model – can be particularly problematic, with such data remaining hidden in the notebooks of experimentalists. This access problem is worsened by publication bias towards successes and weak incentives to reveal data on failures (Raccuglia et al., 2016; Krieger, 2021).

We have noted how the basic exploitation model of Section 2 be extended to allow for the value of exploration. To the extent that the primary output of exploration is data, and at least some of this data spills over to other users, there is also a potential bias towards excessive exploitation of existing data instead of more exploratory research to generate new data. This suggests the value of public subsidies for research that explores the design space. At the firm level, the centrality of data highlights the importance of firms being part of open innovation networks – both providing and gaining access from their active participation in data sharing.²⁷

Third, the evolution of AI as a technology for discovery will depend on how the downstream applications feedback to the development of the AI technology itself. We identify two important dimensions: (i) whether development is demand- or supply driven; and (ii) whether development is task-specific or task-general.

Although economists have generally adopted a positive attitude towards technological change, there is increasing attention to the importance of the *direction* of such change to the ultimate benefits for society (see, e.g., Acemoglu and Johnson, 2023). Science is interesting in this context in that it is affected by AI (as highlighted by our model) and is itself a driver of AI development. To the extent that the societal effects of AI-aided science are positive – new drugs, new treatment targets, new energy-efficient materials, etc. – policies may be required to ensure that the allocation of efforts to develop AI support scientific applications rather than, say, the further development of LLMs.

The debate over the relative importance of supply-driven versus demand-driven development has a long history in the economics of innovation. As AI development becomes increasingly concentrated in large corporations due to scale requirements, policy may need to address the need to invest in science-focused AI research as a public good.

Jacob Schmookler argues strongly for the primacy of demand in his classic early study, *Invention and Economic Growth*:

As different classes of goods become relatively more important than before, the yield to inventive effort in different fields will tend to change correspondingly. And if we further grant that inventive effort is influenced by prospective yield, the direction of inventive effort will shift (Schmookler, 1966, p. 180).

While largely agreeing with Schmookler on the importance of demand, Nathan Rosenberg (1974) advocates a more nuanced position in which both supply (notably advancements in science) and demand (driven by end-user needs) drive the innovation process:

The burden of my argument here is that the allocation of inventive resources has in the past been determined jointly by demand forces which have broadly shaped the shifting payoffs to successful innovation, together with supply-side forces which have determined both the probability of success within any particular timeframe as well as the prospective cost of producing a successful invention (Rosenberg, 1974, p. 103).

While there is little doubt that the initial advances in machine learning – including the development of the foundations for deep learning and self-supervised learning – had their origins in academia, driven by curiosity and the imperative of solving scientific problems, this balance appears to have shifted over time. As already noted, even in its more purely science-driven phase, AI is unusual in its emphasis on developing techniques that work even where the theoretical foundations for why they work lags behind. But, as the extent of real world applications has expanded, firm-specific needs have had an increasing influence on the development trajectory of the technology, much of which now occurs outside of academia or, where academics remain centrally involved, in collaboration with downstream users, often involving joint appointments between universities and large platform firms.

The second useful dimension of development is the move from a set of diverse task-specific models to multi-task models. Using François Chollet’s typology of degrees of generalization (Chollet, 2019, pp. 11–12), we can think of this development as the move from *local* generalization (i.e., “the ability of a system to handle new points from a known distribution for a single task or a well-scoped set of known tasks”) to *broad* generalization (i.e., “the ability of a system to handle a broad category of tasks and environments without further human intervention”). We note that both degrees of generalization fall short of what Chollet calls *human-centric extreme* generalization where “the scope considered is the space of tasks and domains that fit within the human experience.”

The idea of “foundation models” was introduced by the Stanford Institute for Human Centered AI (HAI), which they define as “models trained on broad data (generally using self-supervision at scale) that can be adapted to a wide range of downstream tasks (see, e.g., Bommasani et al., 2021).” These models have been central to recent developments in AI and include such large language models (LLMs) as Google’s PaLM, Anthropic’s Claude, and OpenAI’s GPT-4, ChatGPT, and DALL-E.²⁸ Beyond language, they also exist for other modalities including images, code, proteins, speech, and molecules.²⁹

Notwithstanding their role in the development of AI as a GPT, concerns have been raised about the increasing importance of such broad

²⁶ See Farboodi and Veldkamp (2021) for fascinating treatment of the importance and measurement of data in the modern economy.

²⁷ Of course, one risk with open-sourced data is that malevolent actors could put them to socially damaging ends.

²⁸ Here GPT stands for generative pre-trained transformer rather than general purpose technology.

²⁹ While GPT-4 is often taken as the paradigmatic case of a foundation model, another prominent example of the evolution of learning systems from the more task-specific to the task-general would be the evolution of DeepMind’s celebrated AlphaGo to AlphaGoZero to AlphaZero (with the latter capable of learning any two-player game). DeepMind also recently introduced a multi-task learning model (or “generalist agent”) that it calls Gato (Reed et al., 2022). The model is trained on >600 distinct tasks (including vision, language and robotic control tasks). In testing, the model is found to perform reasonably well in many tasks compared to task-specific model benchmarks.

generalization models. One concern is that as the underlying platforms become increasingly homogenous over time, their (possibly poorly understood) weaknesses as well as their strengths get transmitted to downstream applications, becoming a source of systemic risk as well as benefit. This risk can be heightened by the “emergent” nature of the capabilities of the learning system. This suggests a role for public regulation of AI products in some areas, though not necessarily regulation of AI research itself.

A further concern is that the control over a core part of the AI infrastructure becomes increasingly concentrated in large firms and thus a source of potential market power. Policy makers will need to be alert to anti-competitive practices, including the ability to monopolize data and access to the underlying foundation models.

Our model also highlights how AI can alter the balance between humans, algorithms, and robots in the production of science itself. The extreme form of an autonomous discovery system would seem to leave limited tasks for humans, with algorithms doing the predictions and robots doing the testing. But the effects on the division of labor are likely to be case specific. The demand for some types of scientists – such as those who can combine domain-specific knowledge and machine learning – are likely to rise, while demand for those whose tasks can be automated is likely to fall. This raises questions for universities and government in relation to the design of training – and re-training – programs.

The changing balance between humans and algorithms as this meta-GPT evolves is likely, then, to affect both productivity growth and income distribution. One fear is that AI will prove to be, in the pithy phrase of [Acemoglu and Restrepo \(2019\)](#), a “so-so technology” – with limited productivity effects while potentially causing substantial disruption to labor markets. But as a meta-GPT, AI-aided innovation may also affect a wider range of sectors by altering the knowledge production function of the overall economy. Shifting the knowledge production function makes it less likely that AI will be a so-so technology when used as a *tool* in the innovation process, as opposed to simply being a *product* of that process.

AI’s growing uses as it matures (not all likely to be welcome) could also mean that its impact as a technology for innovation could extend far beyond the sectors traditionally associated with information technology, providing new blueprints for “rearranging atoms” and not just for “rearranging bits.” And, of course, the creative destruction induced by these innovation processes may affect labor demand much more broadly

than the narrow scientist skill markets we have considered here. In any case, new dynamic models of the innovation process that highlight costly search over combinatorial search spaces are needed to help us better understand this important new force that could change the balance between humans and machines at the economy’s innovation frontiers.

What is the main testable implication of the model? It is simple to state but difficult to implement – access to AI-based prediction models will increase scientific discovery and innovation. We refer to this as the *Hassabis hypothesis*. Demis Hassabis, co-founder of DeepMind (now part of Google), has been an evangelist for the potential of AI to speed scientific discovery. He has noted three requirements that make a scientific problem amenable to an AI-aided solution: a combinatorial search space (too large for exhaustive search); a clear objective function for training the prediction model; and sufficient data or capability to simulate that data to train the model. We suggest adding a fourth – poor alternative predictive models for prioritizing search over the design space. When these conditions are present, the Hassabis hypothesis is essentially that the space of amenable problems that cannot be solved by other means is large. AlphaFold can reasonably be seen as a proof of concept – albeit one with significant real-world implications. Time – and empirical research – will tell if it is a fluke or a harbinger of a new era of discovery.

CRedit authorship contribution statement

Ajay Agrawal: Conceptualization, Writing – review & editing. **John McHale:** Conceptualization, Writing – original draft, Writing – review & editing. **Alexander Oettl:** Conceptualization, Writing – review & editing.

Declaration of competing interest

John McHale reports financial support was provided by Science Foundation Ireland. Declaration of Interests: Agrawal: <https://agrawal.ca/disclosure> McHale/Oettl: None.

Data availability

No data was used for the research described in the article.

Appendix A. Appendices

A.1. Generalization of the model to heterogeneous payoffs and costs of testing

The model in [Section 2.1](#) assumes that the payoffs conditional on success and the cost of testing are each the same for all potential design combinations. More specifically, we assume: $v_i = 1, \forall i \in \{1, \dots, 2^N\}$ and $c_i = c, \forall i \in \{1, \dots, 2^N\}$. This allows the ranking of combinations – and thereby the order of search – to depend only on the probabilities of success of those combinations. Here, we generalize the results to the case of heterogeneous payoffs and costs.

As shown in [Weitzman \(1979\)](#), the optimal order of search for a profit-maximizing decision maker can be specified in declining order of “reservation prices.” (See [Agrawal et al., 2023](#), for a simple proof of the Bernoulli case.) The reservation price is the “sure thing” that makes the decision maker (i.e., the innovator) indifferent between a lottery of obtaining a success with probability p_i and otherwise receiving the sure-thing payoff, and the sure-thing payoff. Letting θ_i denote this reservation price, it is therefore determined as:

$$(v_i - c_i) + (1 - q_i)(\theta_i - c_i) = \theta_i$$

$$\Rightarrow \theta_i = v_i - \frac{c_i}{q_i} \quad (\text{A1.1})$$

[Agrawal et al. \(2023\)](#) show that for combinations with a reservation price greater than or equal to zero the optimal search strategy with such heterogeneous combinations (or “Bernoulli boxes”) is to search the combinations in order of reservation prices from highest to lowest (randomizing in the case of equal reservation prices) and to stop the search when a success is found. If no success is found once all the non-negative reservation prices have been searched, then search should stop.

Again, indexing a combination in the ranked list by its rank z , the maximum search that maximizes expected profit, z^* , is given by the largest rank such that $\theta_{z^*} \geq 0$. It is easily verified from A1.1 that this is equivalent to the largest rank such that $q_{z^*} \cdot v_{z^*} - c_{z^*} \geq 0$.

While the ranking is no longer determined by the success probabilities alone, we are still able to define the hazard function based on the new

optimal ranking. However, unlike the monotonically declining success function in the homogenous payoffs and costs case, this success function can rise and fall with z . Notwithstanding this non-monotonicity, we can apply the analysis of Section 2 to this hazard function to determine the expected innovation output, expected search duration, and expected profit. Moreover, in addition to the information on the heterogeneous payoffs and costs, the related cumulative hazard function provides sufficient information to determine these expected innovation outcomes. The corresponding equations to (10) and (11) for expected innovation output and expected search duration, respectively, are identical (though we note that the methods of determining the optimal maximum search length are different). Given the heterogeneous payoffs and costs, expected innovation profit is then,

$$\begin{aligned}\Lambda(z^*) &= \sum_{z=1}^{z^*} S(z-1)(q_z v_z - c_z) \\ &= \sum_{z=1}^{z^*} e^{-H(z-1)}(q_z v_z - c_z)\end{aligned}\quad (\text{A1.2})$$

Once we for heterogeneity across boxes in the value of payoffs, it is easy to show that, all else equal, “riskier” boxes will tend to be prioritized in the search. We assume $c_i = c$ to limit the heterogeneity to the value of boxes conditional on success. Now consider two boxes, Box 1 and Box 2, such that $v_1 > v_2$ and $q_1 < q_2$. Box 1 has a high payoff but a low probability of success; Box 2 has a low payoff but high probability of success. From A1.1 it is clear that Box 1 has a higher reservation price and will be searched first. Therefore, allowing for heterogeneous success outcomes, the innovator will tend to prioritize “riskier” boxes in the search.

A.2. Expected duration of search and the cumulative survival function

Proposition: The expected duration of search with restricted size, z^* , is equal to the sum of the survival probabilities up to $z^* - 1$.

Proof: The expected duration of the restricted search is,

$$D(z^*) = \sum_{z=1}^{z^*} f(z)z + S(z^*)z^* \quad (\text{A2.1})$$

That is, the expected duration of the restricted search is given by the sum of all search durations up to z^* weighted by their probabilities of being the first success plus the maximum duration weighted by the probability that the search will not have achieved a success on completion of the final test.

We now use the fact that the probability of a first success at a given rank, z , is equal to the negative of the change in the survival probability – i.e., $f(z) = -\Delta S(z)$ – to rewrite (A2.1) as,

$$D(z^*) = \sum_{z=1}^{z^*} -\Delta S(z)z + S(z^*)z^* \quad (\text{A2.2})$$

Expanding the first term on the right-hand-side and cancelling yields,

$$D(z^*) = S(0) + S(1) + \dots + S(z-1) - S(z^*)z^* + S(z^*)z^* = \sum_{z=1}^{z^*} S(z-1) \quad (\text{A1.3})$$

It is useful to illustrate with an example taking $z^* = 3$. Expanding (A2.2) yields:

$$\begin{aligned}D(3) &= (S(0) - S(1))1 + (S(1) - S(2))2 + (S(2) - S(3))3 + S(3)3 \\ &= S(0) + S(1) + S(2)\end{aligned}$$

References

- Acemoglu, D., Johnson, S., 2023. *Power and Progress: Our Thousand-Year Struggle over Technology and Prosperity*. Basic Books, London.
- Acemoglu, D., Restrepo, P., 2019. Automation and new tasks: how technology displaces and reinstates labor. *J. Econ. Perspect.* 33 (2), 3–30.
- Aghion, P., Howitt, P., 1992. A model of growth through creative destruction. *Econometrica* 60, 323–351.
- Agrawal, A., McHale, J., Oettl, A., 2019. Finding needles in haystacks: Artificial intelligence and recombinant growth. In: Agrawal, Gans, Goldfarb (Eds.), *The Economics of Artificial Intelligence: An Agenda*. University of Chicago Press, Chicago, IL.
- Agrawal, A., Gans, J., Goldfarb, A., 2022. *Power and Prediction: The Disruptive Economics of Artificial Intelligence*. Harvard Business Press, Boston.
- Agrawal, A., McHale, J., Oettl, A., 2023. Superhuman science: How artificial intelligence may impact innovation. *J. Evol. Econ.* 33, 1473–1517.
- Arthur, B.W., 2009. *The Nature of Technology: What it Is and how it Evolves*. Penguin Books, London.
- Aspuru-Guzik, A., Persson, K., 2018. *Materials Acceleration Platform: Accelerating Advanced Energy Materials Discovery by Integrating High-Throughput Methods and Artificial Intelligence*. Mission Innovation.
- Bianchini, S., Müller, M., Pelletier, P., 2022. Artificial intelligence in science: an emerging general method of invention. *Res. Policy* 51 (10), 1–15.
- Bloom, N., Jones, C.I., Van Reenen, J., Webb, M., 2020. Are ideas getting harder to find? *Am. Econ. Rev.* 110 (4), 1104–1144.
- Bommasani, R., Hudson, D.A., Adeli, E., Altman, R., Arora, S., von Arx, S., Liang, P., 2021. On the opportunities and risks of foundation models. *arXiv preprint. arXiv:2108.07258*.
- Bresnahan, T., 2010. General purpose technologies. *Handbook of the Economics of Innovation* 2, 761–791.
- Bresnahan, T.F., Trajtenberg, M., 1995. General purpose technologies ‘engines of growth’? *J. Econ.* 65 (1), 83–108.
- Brynjolfsson, E., Rock, D., Syverson, C., 2019. Artificial intelligence and the modern productivity paradox: A clash of expectations and statistics. In: *The Economics of Artificial Intelligence: An Agenda*. University of Chicago Press, pp. 23–57.
- Brynjolfsson, E., Rock, D., Syverson, C., 2021. The productivity J-curve: how intangibles complement general purpose technologies. *Am. Econ. J. Macroecon.* 13 (1), 333–372.
- Cavalli, G., 2023. *How Scientific Organizations Adapt to Methodological Advances in Artificial Intelligence: The Impact of AlphaFold1*. Working Paper. Rotman School of Management, University of Toronto.
- Chollet, F., 2019. On the measure of intelligence. *arXiv preprint. arXiv:1911.01547*.
- Clancy, M.S., 2018. Inventing by combining pre-existing technologies: patent evidence on learning and fishing out. *Res. Policy* 47 (1), 252–265.
- Cockburn, I.M., Henderson, R., Stern, S., 2019. 4. The impact of artificial intelligence on innovation: An exploratory analysis. In: *The Economics of Artificial Intelligence*. University of Chicago Press, pp. 115–148.
- Crafts, N., 2021. Artificial intelligence as a general-purpose technology: an historical perspective. *Oxf. Rev. Econ. Policy* 37 (3), 521–536.
- Dasgupta, P., Stiglitz, J., 1980. Uncertainty, industrial structure, and the speed of R&D. *Bell J. Econ.* 11 (1), 1–28.

- David, P.A., 1990. The dynamo and the computer: an historical perspective on the modern productivity paradox. *Am. Econ. Rev.* 80, 355–361.
- Farboodi, M., Veldkamp, L., 2021. A model of the data economy. NBER Working Paper, 28427.
- Fleming, L., Sorenson, O., 2004. Science as a map in technological search. *Strateg. Manag. J.* 25, 909–928.
- Griliches, Z., 1957. Hybrid corn: an exploration in the economics of technological change. *Econometrica* 501–522.
- Hassabis, D., 2022. Using AI to accelerate scientific discovery. In: Speech at the Institute for Ethics in AI, University of Oxford. Speech available at: (441) Dr Demis Hassabis: [Using AI to Accelerate Scientific Discovery - YouTube](#).
- Hötte, K., Tarannum, T., Verendel, V., Bennett, L., 2022. AI technological trajectories in patent data: general purpose technology and concentration of actors. In: INET Oxford Working Paper No. 2023-09.
- Jones, C.I., 2021. Recipes and Economic Growth: A Combinatorial March down an Exponential Tail, w28340. National Bureau of Economic Research.
- Jovanovic, B., Rousseau, P.L., 2005. General purpose technologies. In: *Handbook of Economic Growth*, 1. Elsevier, pp. 1181–1224.
- Jumper, J., et al., 2021. Highly accurate protein structure prediction with AlphaFold. *Nature* 596, 583–589.
- Kortum, S., 1997. Research, patenting, and technological change. *Econometrica* 65 (6), 1389–1419.
- Krenn, et al., 2022. On scientific understanding with artificial intelligence. *Nature Reviews Physics* 4, 761–769.
- Krieger, J., 2021. Trials and terminations: learning from competitor's R&D failures. *Manag. Sci.* 67 (9), 5525–5548.
- Ludwig, J., Mullainathan, S., 2023. Machine learning as a tool for hypothesis generation. In: NBER Working Paper 31017. MA, Cambridge.
- Mokyr, J., 2002. *The Gifts of Athena: Historical Origins of the Knowledge Economy*. Princeton University Press, Princeton.
- Nelson, R.R., Winter, S.G., 1982. *An Evolutionary Theory of Economic Change*. Harvard University Press, Cambridge, MA.
- Pyzer-Knapp, E.O., Li, K., Aspuru-Guzik, A., 2015. Learning from the Harvard clean energy project: the use of neural networks to accelerate materials discovery. *Adv. Funct. Mater.* 25, 6495–6502.
- Raccuglia, P., Elbert, K.C., Adler, P.D., Falk, C., Wenny, A.B., Mollo, A., Zeller, M., Fridler, S.A., Schrier, J., Norquist, A.J., 2016. Machine-learning-assisted materials discovery using failed experiment. *Nature* 533 (7601), 73–76.
- Raghu, M., Schmidt, E., 2020. A Survey of deep learning for scientific discovery arXiv: 2002.11755v1.
- Ramsundar, B., Eastman, P., Walters, P., Pande, V., 2019. *Deep Learning for the Life Sciences: Applying Deep Learning to Genomics, Microscopy, Drug Discovery & More*. O'Reilly Media, Sebastopol, CA.
- Reed, S., Zolna, K., Parisotto, E., Colmenarejo, S.G., Novikov, A., Barth-Maron, G., de Freitas, N., 2022. A generalist agent. arXiv preprint. arXiv:2205.06175.
- Romer, P., 2008. Economic growth, published in *The Concise Encyclopaedia of Economics*, Library of Economic Liberty. available at: <http://www.econlib.org/library/Enc/EconomicGrowth.html>.
- Rosenberg, N., 1974. Science, invention and economic growth. *Econ. J.* 84 (333), 90–108.
- Rosenberg, N., 1994. *Inside the Black Box: Technology, Economics, and History*. Cambridge University Press, Cambridge, U.K.
- Rosenberg, N., 1998. Chemical engineering as a general purpose technology. In: Helpman, Elhanan (Ed.), *General Purpose Technologies and Economic Growth*. MIT Press, Cambridge, MA.
- Schmookler, J., 1966. Invention and economic growth. In: *Invention and Economic Growth*. Harvard University Press.
- Schumpeter, J.A., 1939. *Business Cycles*, 1. McGraw-Hill New York.
- Tabor, D.P., Roch, L.M., Saikin, S.K., Kreisbeck, C., Sheberla, D., Montoya, J.H., Dwaraknath, S., Aykol, M., Ortiz, C., Tribukait, H., et al., 2018. Accelerating the discovery of materials for clean energy in the era of smart automation. *Nat. Rev. Mater.* 3, 5–20.
- Usher, A.P., 1929. *A History of Mechanical Inventions*, revised edition. McGraw-Hill, New York.
- Weitzman, M.L., 1979. Optimal search for the best alternative. *Econometrica* 641–654.
- Weitzman, M.L., 1998. Recombinant growth. *Q. J. Econ.* 113, 331–360.