

Combinatorial innovation, evidence from patent data, and mandated innovation

by

Matthew Clancy

A dissertation submitted to the graduate faculty
in partial fulfillment of the requirements for the degree of

DOCTOR OF PHILOSOPHY

Major: Economics

Program of Study Committee:
GianCarlo Moschini, Major Professor
Gray Calhoun
David Frankel
Harvey Lapan
Sergio Lence

Iowa State University

Ames, Iowa

2015

TABLE OF CONTENTS

	Page
ACKNOWLEDGMENTS	iv
CHAPTER 1. INTRODUCTION	1
CHAPTER 2. COMBINATORIAL INNOVATION: EVIDENCE FROM TWO CENTURIES OF PATENT DATA	
Abstract.....	4
0. Introduction	4
1. Background	7
2. Model Basics	10
3. The Firm's Problem	14
4. Predictions	16
5. Data	18
6. Probability of Use.....	25
7. Patents Issued Per Year.....	33
8. Discussion	45
9. Conclusions	46
References	47
Appendix – Robustness Checks for Section 7	49
CHAPTER 3. OPTIMAL RESEARCH STRATEGIES WHEN INNOVATION IS COMBINATORIAL	
Abstract.....	54
0. Introduction	54
1. Background	56
2. Model Basics	58
3. Special Case 1: Affinity is Known.....	63
4. Research and Learning.....	64
5. Special Case 2: Research as a Multi-armed Bandit Problem	67
6. The General Case	73
7. Numerical Simulations.....	76
8. Numerical Analysis of the General Case	81
9. Conclusions	92
References	93
Appendix A1: Gittins Index	95
Appendix A2: Program Details	97

CHAPTER 4. MANDATES AND THE INCENTIVES FOR ENVIRONMENTAL INNOVATION

Abstract.....	102
1. Background	102
2. The Model	107
3. Mandate vs. Carbon Taxes: A Comparison.....	118
4. Multiple Innovators.....	119
5. Mandate vs. Carbon Tax with Multiple Innovators	126
6. Some Numerical Results.....	129
7. Conclusion.....	139
References	140
CHAPTER 5. CONCLUSION.....	143

ACKNOWLEDGMENTS

I would like to thank my major professor, GianCarlo Moschini, for years of patient feedback and advice. I am also grateful for three years of support from the USDA National Needs Fellowship, which gave me the time and freedom to develop these ideas. I owe much to the members of the Iowa State University department of economics faculty that helped facilitate this support. I would also like to specifically thank the rest of committee, Gray Calhoun, David Frankel, Harvey Lapan, and Sergio Lence for helpful comments and criticism.

CHAPTER 2

COMBINATORIAL INNOVATION: EVIDENCE FROM TWO CENTURIES OF PATENT DATA

Matthew S. Clancy

Abstract

This paper presents an original model of knowledge production, and tests several predictions of the model using a novel dataset built from 8.3 million US patents. In this model, new ideas are built by combining pre-existing technological building blocks into new combinations. The outcome of research is always stochastic, but firms are Bayesians who learn which sets of technological building blocks tend to yield useful discoveries and which do not. Consistent with this model's prediction, I show that the number of patents granted in a particular technology class increases in the years after new useful combinations of technology first appear in the class. Moreover, after new combinations first appear, I show subsequent patents are more likely to draw on the same combination of technology, consistent with firms learning the technologies can be fruitfully combined. Patents are also more likely to combine technologies that have already been combined with many of the same (other) technologies, even if they have never been combined with each other. Finally, I show that the probability of using a combination declines over time, and that the total number of patents granted in a technology class also declines over time, if there are not new connections between technologies continuously discovered. This is consistent with the model's predictions about firms using up all the useful ideas that can be built from a fixed set of technological building blocks.

0. Introduction

This is a paper about where ideas come from. There is a large and fruitful literature on the economics of innovation, but the nature by which ideas are found is largely treated as a black box. As Weitzman (1998) expressed it “[In endogenous growth models] “New ideas” are simply taken to be some exogenously determined function of “research effort” in the spirit of a humdrum conventional relationship between inputs and outputs.” We may add to Weitzman’s “research effort” additional inputs such as human capital or measures of knowledge itself (e.g., the number of

varieties of machines, or their quality), but the birth of ideas is usually not given strong microfoundations.

Of course, given the fecundity of the literature, this approach is justified in many applications. Nevertheless, there are questions that are difficult to answer without a better theory of how ideas are created. Gordon (2012), for example, argues the supply of good ideas has run out, a theme also discussed in Cowen (2011). Conversely, Brynjolfsson and McAfee (2014), expressing a view common to technology companies today, argue we are entering a period of brilliant technological progress. We need a theory of where ideas come from to approach this controversy. Such a theory would also be useful in addressing questions related to science policy: tightening budgets increasingly require agencies to make difficult decisions in funding different branches of knowledge, at different stages of exploratory and applied research.

This paper proposes a microfounded model of innovation, and tests some predictions of the model using with a novel dataset constructed from US patent data. It models the creative act as drawing on a pre-existing set of technological building blocks. These “components” must be assembled into a novel combination by an innovating firm to develop a new technology. The way these components interact with each other determines whether or not the new technology is useful. Moreover, I explicitly model the way firms *learn* how components tend to interact.

Consider the internal combustion engine as an illustrative example. The internal combustion engine is a single idea created by inventors, but is also a combination of pistons, crankshafts, flywheels, valves, and combustible chemicals.¹ All of these ingredients existed in some shape or another, before the invention of the internal combustion, but the internal combustion engine did not exist until they were brought together. At the same time, an internal combustion engine is more than the combination of these building blocks: piling pistons, crankshafts, flywheels, valves, and combustible chemicals in a heap, for example, would not yield up an engine. Instead, the connections between the elements matters, with all the elements working in harmony, rather than in opposition, to perform useful tasks. The combustible chemicals drive the pistons, the crankshaft converts this back and forth motion into jerky rotational motion, and the flywheel converts the crankshaft’s jerky motion into comparatively smooth rotational motion, and so on.

When conducting the research and design (R&D) that would eventually culminate in an internal combustion engine researchers began with some information. They knew how the components they

¹ The following example draws on Dartnell (2014), pgs. 201-206.

intended to use were likely to interact, because the components had been used before in other contexts. Waterwheels also use crankshafts and pistons, and a potter's wheel couples a flywheel to a jerky source of motion. The likely interaction of these pairs of technologies (crankshaft and piston, flywheel and piston) could be inferred from these other contexts. In this framework for thinking about innovation, knowledge spills over in both the recycling of component parts, and in the understanding of how parts may be used together.

This framework yields a number of predictions. For example, firms will be more likely to combine building blocks that they know tend to work well together, or at least belong to a set of components that firms know to be mutually compatible. Moreover, the more technological building blocks that researchers know can be fruitfully combined with each other, the more ideas that will be worth developing. At the same time, the number of combinations that can be built from a finite set of building blocks is also finite, and over time, firms will run out of things to build.

I test these predictions using a novel dataset on US patents. I exploit the US Patent Office's Classification system to assign each patent a set of technological building blocks that comprise it. Each such building block is proxied by one of the patent office's technology *subclasses*, a designation considerably more fine-grained than the more typical technology classes that are used in most empirical patent papers. I show that aggregate patenting activity rises in the year after patents successfully combine subclasses that are otherwise rarely used together, and that new patents are more likely to draw on these newly combined classes. Moreover, I show that, in the absence of these new connections, patenting activity falls off. This is consistent with the predictions of my model.

I begin by providing some further background (section 1) before a formal model of combinatorial knowledge production is introduced (section 2). I then set up a model of a single inventor's problem (section 3). To motivate my empirical application, I make a series of predictions when there are multiple firms operating within the same technological sector (section 4). This model predicts the number of ideas produced declines over time, unless useful connections between elements are continually discovered. It also predicts new connections will be made where they are most complementary to existing connections, and that the relationship between a connection's past and future use is increasing. I construct a novel dataset from US patent data (section 5), which I use to test these predictions (sections 6 and 7). After some discussion of the results (section 8), I consider some directions for future research (section 9).

1. Background

The internal combustion is just one example of the ways technology and ideas themselves can be viewed as fundamentally combinatorial. Any physical technology can be broken down into component parts, while invented procedures can be broken down into steps and actions. Non-physical creations can also be understood as combinations. Works of fiction draw on a common set of themes, styles, character archetypes, and other tropes; musical compositions rely on combinations of instruments, playing styles, and other conventions; and paintings deploy common techniques, symbols, and conventions. Indeed, even abstract ideas can be understood as combinations of concepts, arguments, mathematical tools, facts, and so forth.

Weitzman (1998) is the first to incorporate this feature of knowledge creation into the economist's knowledge production function. In Weitzman's model, innovation consists of pairing "idea-cultivars"² to see if they yield a fruitful innovation (a new idea-cultivar), where the probability an idea-pair will bear fruit is an increasing function of research effort. If successful, the new idea-cultivar is included in the set of possible idea-cultivars that can be paired in the next period. Weitzman's main contribution is to show that combinatorial processes eventually grow at a rate faster than exponential growth processes, so that, absent some extreme assumptions about the cost of research, in the limit growth eventually becomes constrained by the share of income devoted to R&D rather than the supply of ideas. Simply put, combinatorial processes are so fecund that we will never run out of ideas, only the time needed to explore them all.

Weitzman's model is echoed in Arthur (2009), who views all technologies as hierarchical combinations of sub-components. Arthur agrees that the internal combustion engine is composed of pistons, crankshafts, flywheels, and so on, but goes further, pointing out that, say, the piston, is itself a combination of two shaped metal components, as well as lubrication, and so forth. These sub-components themselves are combinations of still further subcomponents (metal alloys, for instance). Ridley (2010) also proposes a model akin to Weitzman's, arguing the best innovations emerge when "ideas have sex."

These models display *recombinant* growth, wherein new combinations subsequently become the raw ingredients of future combinations. In contrast, this model is concerned primarily with the ways a *fixed* set of components can be reconfigured in many ways, and the interaction between two elements is used again and again to achieve different purposes. I have no doubt that Weitzman's

² So-called because the hybridization of ideas in the model is analogous to the hybridization of plant cultivars.

model is appropriate for understanding where the elements of combination themselves come from, in the long run. This model is concerned with the shorter run, where the set of elements available for combination is relatively static and new elements are rare and unanticipated shocks to the set of available technological building blocks.

Several other papers have modeled the innovation process as one of learning about an underlying space of potential ideas. Jovanovic and Rob (1990) represents a technology by an infinite vector, each element of which ranges between 0 and 1. Technologies are production functions and agents learn the mapping from technology vectors to productivity via Bayesian updating. Research consists in changing the values of the elements in a vector and observing the labor productivity associated with the new vector. Kauffman, Lobo and Macready (2000) and Auerswald et al. (2000) follow Jovanovic and Rob (1990) in thinking of technologies as a large combination of distinct operations, although here the length of a technology vector is finite and each element can take on one of a finite number of states (rather than ranging over a continuous interval). The mapping between each technology vector and its productivity level is called a fitness landscape. When states are interdependent, the authors show this landscape is characterized by many local maxima. Innovation in such a model consists of exploring the fitness landscape by changing different operations.

The model presented in this paper differs from the above papers by modeling the learning process as focused on the relationships between components in an idea directly. To illustrate the difference, consider again the internal combustion engine. In my model, I asserted that an inventor would know how a piston and crankshaft were likely to interact in an internal combustion engine by observing how they interact in a waterwheel. Although the engine and the waterwheel are otherwise very different, useful information can be extracted from the waterwheel that can be applied to the engine. In the models of Kauffman, Lobo and Macready (2000) or Jovanovic and Rob (1990), the usefulness of the waterwheel would instead be defined by the technological distance between the waterwheel and the engine, which is likely to be quite large. This paper's model implicitly assumes the most important characteristics of an idea can be decomposed into a function of the set of pair-wise interactions between all its building blocks. The other papers instead consider the entire idea as a fixed point, and new possible ideas that are near the fixed point will have performance more highly correlated with it. This emphasis on the decomposability of technology into pair-wise interactions greatly simplifies empirical analysis.

Another line of literature uses a combinatorial framework in empirical applications. Such papers invariably use patents or academic papers as measures of innovation. There tend to be two ways of measuring the “components” of an idea. One approach is to use the citations of a patent or paper, generally grouped into technology or discipline categories, and another is to look at the technology or discipline categories assigned directly to a patent or paper, as this paper does. Meanwhile, a patent or paper’s value tends to be measured by the number of citations it has received, in line with work by Trajtenberg (2002) and Harhoff et al. (1999) that shows such measures are correlated with independent measures of patent value. Studies using a combinatorial framework usually attempt to predict the value of patents (or papers) by using traits of the combination, such as the extent of unusual combinations or frequency with which ideas have been used. As far as I have been able to determine, this is the first paper that attempts to predict changes in aggregate patenting activity by employing a combinatorial innovation framework.

The paper most similar in spirit to this one is Fleming (2001), which examines a sample of 17,264 patents through a similar combinatorial lens. Fleming (2001) exploits the fact that most patents are assigned to more than one technological subclass, interpreting each sub-class assignment as a component. Fleming then simply counts the number of times such a combination of sub-classes has resulted in a patent, finding that less commonly used combinations obtain more citations. Fleming also computes a measure of combination familiarity, which is a weighted sum of the number of times a combination of sub-classes has been patented, with more recent combinations weighted more heavily. This familiarity metric, which Fleming interprets as accounting for the distance of search, is also associated with more citations. Other papers using patent data include Nemet (2012), Nemet and Johnson (2012), and Shoenmakers (2010). These papers all rely on citation data to proxy for the technological building blocks used by an idea, and tend to show unusual connections between technological fields in citations results in a higher citation rate. Schilling (2011), obtains a similar result for a sample of academic papers.

This paper differs from these other empirical papers in a number of respects. First, my dataset is considerably larger than previous studies, encompassing all 8.3 million US utility patents granted through 2012. Second, this paper does not use patent citation data. To measure the technological building blocks that comprise an idea, it uses patent subclasses, albeit in a novel way (discussed in section 5). One virtue of this approach is that a patent’s technology classifications are supplied by a single ostensibly neutral arbiter, namely the patent office, while citations are made by both applicants and the patent office. Third, this paper develops an explicit theoretical model, from which

predictions are derived (sections 2-4). These predictions do not require me to have any theory of citations. Instead, they pertain to the *number* of patents that will be granted in future period, as well as the probability a pair of technological components will be assigned to a patent. As far as I have been able to determine, mine is the first paper to show combinatorial factors have a measurable impact on the aggregate rate of future patent activity. I turn now to the presentation of this model.

2. Model Basics

2.1. The Knowledge Production Function

I now describe formally how ideas are created in this model.

Definition 1: Primitive Elements. Let Q denote the set of primitive elements of knowledge q that can be combined with other elements to produce ideas, where $q \in Q$.

Definition 2: Pairs. Let p denote a two-element subset of Q , or “pair,” and P denote the set of two-element subsets of Q , where $p \in P$.

Definition 3: Ideas. An *idea* d is a set of pairs p , satisfying the condition that if $p_0 \in d$ and $p_1 \in d$, then $p \in d$ for any $p \subseteq p_0 \cup p_1$.

A fixed set Q of technological building blocks can be assembled into ideas, where any idea must contain at least two q from the set Q . For convenience, I define ideas in terms of the *pairs* of elements contained therein. For example, an idea combining elements q_1 , q_2 , and q_3 is represented as the set of subsets $((q_1, q_2), (q_1, q_3), (q_2, q_3))$. The condition attached to Definition 3 merely insures the pairs between all elements in the idea are included in the idea, so that we do not have ideas such as $((q_1, q_2), (q_1, q_3))$, which uses elements q_1 , q_2 , and q_3 but does not include the pair corresponding to (q_2, q_3) .

There are three important concepts in this model.

Definition 4: Compatibility. The *compatibility* of pair p in idea d is $c(p, d) \in \{0, 1\}$. When $c(p, d) = 1$ then the pair p is compatible in d . When $c(p, d) = 0$ then the pair p is incompatible in d .

Note that $c(p, d) = c(p, d')$ is not generally true. The compatibility of a pair may be equal to 1 in one idea and 0 in another.

Definition 5: Affinity. The probability a pair p is compatible defines its *affinity* $a(p) \in [0, 1]$.

The notions of compatibility and affinity are related as follows:

$$c(p, d) = \begin{cases} 1 & \text{with probability } a(p) \\ 0 & \text{with probability } 1 - a(p) \end{cases} \quad (1)$$

Essentially, this model assumes pairs of elements have an underlying tendency to be compatible or incompatible, and this tendency is described by the affinity of the pair.

Lastly, ideas are either *effective* or *ineffective*, where an idea is effective if and only if all the pairs of its constituent elements are compatible.

Definition 6: Efficacy. An idea d is *effective*, represented by $e(d) = 1$, iff $c(p, d) = 1 \forall p \in d$. In all other cases, represented by $e(d) = 0$, idea d is *ineffective*.

Restated, affinity determines the probability a pair is compatible, and when all pairs in an idea are compatible, the idea is effective. We may imagine ideas as sets of interacting elements that must be mutually compatible for the idea to prove useful. If any two elements are incompatible, I assume the idea suffers a catastrophic failure that renders it unfit for use. Note the probability an idea is effective can be written as:

$$\Pr(e(d) = 1) = E[e(d)] = \prod_{p \in d} a(p) \quad (2)$$

This is the joint probability that every pair in the idea is compatible. Ideas are most likely to be effective when they are composed exclusively of elements that have a high affinity for each other, and least likely to be effective when composed of elements with a low affinity for each other.

2.2. Beliefs

The affinity $a(p)$ between a pair of elements is *ex ante* unknown to researchers. Instead, researchers are Bayesians with prior beliefs over the possible distribution of $a(p)$. Though the researcher does not observe $a(p)$ directly, she can make educated guesses based on the tendency of p to be compatible or incompatible. Using her beliefs about the affinities between all pairs in an idea, she can compute the probability an idea will be effective. In more formal terms, a crucial part of the discovery process is the inference of likely affinity values from the compatibility or incompatibility of component-pair interactions.

I impose one assumption on the researcher's beliefs:

Assumption 1: Independence of Affinity. The researcher believes $a(p)$ is independently distributed for all p .

As long as this assumption stands, the updating of beliefs about any $a(p)$ depends only on observations on the pair p alone. If $a(p)$ were not independently distributed, it would be necessary to also take into account the observations on correlated pairs, greatly complicating the problem.

Each observation of compatibility is the outcome of a Bernoulli trial governed by the pair's true affinity, with the two possible states being compatibility (probability $a(p)$) or incompatibility (probability $1 - a(p)$). Given s instances of compatibility ("success") and f instance of incompatibility ("failure"), the researcher updates her beliefs according to Bayes law under the Bernoulli distribution:

$$\Pr(a(p) = \tilde{a} \mid s, f) = \binom{s+f}{s} \tilde{a}^s (1 - \tilde{a})^f \frac{\Pr(a(p) = \tilde{a})}{\int_0^1 \binom{s+f}{s} a^s (1 - a)^f \Pr(a(p) = a) da} \quad (3)$$

where $\binom{s+f}{s} \tilde{a}^s (1 - \tilde{a})^f = \Pr(s, f \mid a(p) = \tilde{a})$.

The expected value of $a(p)$ is given by:

$$E[a(p) | s, f] = \frac{\int_0^1 a^{s+1} (1-a)^f \Pr(a(p) = a) da}{\int_0^1 a^s (1-a)^f \Pr(a(p) = a) da} \quad (4)$$

Equation (4) says the expected value of $a(p)$ is an integral over all possible values of $a(p)$, where every a is weighted by $a^s (1-a)^f$ and a factor that normalizes the sum of probabilities to 1. As s increases, the relative weight attached to higher values of a increases more than the weight attached to low values, and $E[a(p) | s, f]$ increases. In the limit, $E[a(p) | s, f]$ converges to 1, as $a^s (1-a)^f$ goes to zero for all $a \neq 1$. Conversely, $E[a(p) | s, f]$ decreases as f increases, and converges to 0 as f grows large relative to s . In other words, researchers believe a pair is more likely to be compatible in the future if it has been compatible in the past, and vice-versa.

The term $a^s (1-a)^f$ can also be written as $\{a^x (1-a)^{1-x}\}^n$ where $x = s/n$ and $n = s + f$. The term $a^x (1-a)^{1-x}$ attains its maximum when $a = x$, so that as n increases, the expected value of $a(p)$ converges to s/n . In other words, as the number of observations increases, researchers come to believe the affinity of a pair is just equal to the proportion of times the pair has been observed compatible.

Given equation (2) and Assumption 1, the expected efficacy of an idea when researcher beliefs are uncertain can be written as:

$$E[e(d)] = \prod_{p \in d} E[a(p) | s(p), f(p)] \quad (5)$$

where $s(p)$ and $f(p)$ denote the number of observations of pair p 's compatibility or incompatibility respectively. This expression captures the core notion of this model: ideas are more likely to be effective ("useful") when composed of elements that have worked well together frequently in the past.

3. The Firm's Problem

Suppose this production function is used by a firm trying to discover effective ideas. The firm knows every element in set Q , and in each period may choose to conduct a research project on some idea d built from the elements in Q .

Definition 7: Possible Ideas. The set D_P is the set of all possible ideas that can be made from elements in Q . It contains all subsets of P that satisfy the condition in Definition 3.

Definition 8: Eligible Ideas. A set of eligible ideas $\tilde{D}$ is a subset of D_P . It is only sensible to conduct research projects on eligible ideas, and when a research project is attempted, the idea is removed from $\tilde{D}$ at the end of the period.

The set $\tilde{D}$ is primarily intended to indicate the set of *untried* ideas, and so it shrinks as research proceeds. I add to this set an additional element, the null set $d_0 \equiv \{\emptyset\}$, which represents the option not to conduct research in a period.

Definition 9: Available Actions. The firm's set of available actions is $D \equiv d_0 \cup \tilde{D}$.

Note that because $d_0 \notin \tilde{D}$, if the firm chooses not to conduct research, then this option is not removed from its action set in the next period.

In principle, the firm “knows” every idea that can be built from elements in Q , in the same sense that I “know” every economics article that can be written with words and symbols in my repertoire. However, just as I do not know whether any of these articles are good until I think more about them, or actually write them out, the firm does not learn if an idea is effective until it decides to conduct research on it.³ Indeed, research is costly, requiring investments of time and other resources. I assume that research on any idea has cost $k(d)$, known to the firm, and that the option d_0 , to do nothing, has $k(d_0) = 0$.

³ Jorge Luis Borges tells a parable of an infinite library containing books with every combination of letter and punctuation mark. In this library, there is a book resolving the basic mysteries of humanity, since every possible book exists, but finding the book and verifying it is true amongst all the gibberish and babel is a daunting task for the library inhabitants. See Borges (1962).

The reward from conducting research may take the form of a patent that pays $\pi(d)$ to the patent holder at the end of the period. Patents may only be obtained for ideas that are eligible and which have been shown to be effective (patents are only issued for useful inventions). Because each chosen idea is removed from the set of eligible ideas D at the end of a period, firms cannot patent the same idea multiple times. I assume the patent value of the outside option d_0 is always zero.

Hence, a firm that chooses to conduct research on idea d expects to receive a net value of:

$$\pi(d)E[e(d)] - k(d) \quad (6)$$

The idea is successful with probability equal to the expected efficacy $E[e(d)]$, in which case the researcher obtains the patent value $\pi(d)$. Whether the idea succeeds or not, the researcher pays up front research costs $k(d)$. This formulation of the innovator's problem is not unusual, except for the term $E[e(d)]$, which is determined by the knowledge production function described earlier and the beliefs of the researcher.

I summarize the firm's information about the prior number of compatibilities and incompatibilities of each pair by the vector I where:

$$I = ((s(p_0), f(p_0)), \dots, (s(p_i), f(p_i)), \dots, (s(p_n), f(p_n))) \quad (7)$$

and where n is the number of pairs in P . I now assume that research on an idea d also reveals information on which pairs are compatible and which are not, in the form of the vector $\omega(d)$.⁴ This vector has the same number of elements as B and is defined so that after conducting a research project on d , the updated information vector I' is given by:

$$I' = I + \omega(d) \quad (8)$$

For example, if a research project on some d' reveals pair p_0 is compatible and pair p_n is incompatible, and does not reveal any other information, then $\omega(d')$ takes the form:

⁴ It may seem to be a natural assumption that research on d reveals the compatibility or incompatibility of every $p \in d$. I do not believe research is so straightforward. For the purposes of this paper, this assumption simplifies the exposition, without much cost. But see Clancy (2015) for a more complete discussion of these issues.

$$\omega(d') = ((1,0), (0,0), \dots, (0,0), (0,1)) \quad (9)$$

Of course, $\omega(d_0) = ((0,0), \dots, (0,0))$ by assumption: agents learn nothing when they choose not to do a research project.

4. Predictions

Consider a technology sector composed of many firms, all engaged in R&D that draws upon a common pool of elements Q . Suppose these firms behave more or less myopically.⁵

A research project is worth pursuing so long as its expected value is positive:

$$E[e(d)]\pi(d) - k(d) \geq 0 \quad (10)$$

This can be rewritten as:

$$\prod_{p \in d} E[a(p)] \geq k(d) / \pi(d) \quad (11)$$

We do not typically observe the right-hand side of equation (11), but I assume it is known to firms. Equation (11) implies that any given idea is more likely to be pursued when the left-hand side is larger, or when the idea's expected affinities are higher. More broadly, within an industry engaged in R&D that draws upon a common pool of elements Q , there will be more ideas worth pursuing, and therefore more patents granted, when the expected affinity of pairs in P_Q is high. This is our first prediction:

Prediction 1: Number of Grants and Affinity. The annual number of patents granted in a technology sector is positively correlated with the average expected affinity of pairs used by the sector.

This prediction will hold in any given year and does not require that the set of elements Q be fixed over time. Indeed, it seems likely to me that new elements are occasionally added to the set of

⁵ In addition to the usual factors that may induce competitive firms to focus on the short-term, innovating firms will behave myopically if knowledge rapidly spills over. When this is the case, competitors can exploit research that yields a return in future periods (for example, by enabling firms to learn currently unprofitable projects are in fact profitable in expectation).

technological building blocks used by a sector.⁶ Over time, however, all ideas satisfying equation (11) are used up (since ideas can only be patented once). This is our second prediction:

Prediction 2: Number of Grants and Time. The annual number of patents granted in a technology sector is negatively correlated with time.

Together, these two predictions imply the stock of available ideas is subjected to two opposing forces. First, a learning effect has the tendency to expand the set of research projects that are *ex ante* profitable to pursue. Second, a fishing out effect uses up profitable ideas.

Predictions 1 and 2 correspond to the aggregate activities of a technology sector, but predictions about individual pairs of elements are also possible. Define the term “use” as follows:

Definition 10: Pair Use. A pair p is used in year t if there is any research project d attempted in year t such that $p \in d$.

In other words, a pair has been used if any research project is attempted on an idea that includes the pair as one of its constituents. For some idea that includes the pair p' , equation (11) can be rearranged to yield:

$$E[a(p')] \geq \frac{k(d) / \pi(d)}{\prod_{p \in d \setminus p'} E[a(p)]} \quad (12)$$

Just as an idea is worth attempting if it satisfies equation (11), the pair p' is worth *using* if there is some eligible idea that satisfies equation (12). Because this condition is more likely to be satisfied when the left-hand side of equation (12) is large, I can make the following prediction.

Prediction 3: Expected Affinity and Use. In any given year, the higher the expected affinity of a pair, the more likely it is to be used.

As discussed above, however, over time research projects satisfying equation (12) are used up.

Prediction 4: Time and Use. The probability a pair will be used decreases over time.

⁶ This could be due to recombinant growth as in Weitzman (1998), or by the discovery of entirely new physical processes that can be co-opted, as in Arthur (2009).

The numerator of equation (12) also highlights the complementarity between pair affinities. If $\prod_{p \in d \setminus p'} E[a(p)] = 0$ then there is no $E[a(p')]$ sufficiently high for the research project to be profitable in expectation. More generally, the higher is $\prod_{p \in d \setminus p'} E[a(p)]$, the more likely equation (12) is true. Furthermore, for the pair p' to be used, there just needs to be one idea that satisfies (12).

Prediction 5: Complementarity and Use. The more pairs with high expected affinity that can be used in conjunction with some pair p' to form eligible ideas, the more likely the pair p' is to be used.

I now turn to the testing of these predictions with data on patents.

5. Data

5.1. Patent Data

To assess these predictions, I use data on patents granted by the US Patent and Trademark Office (USPTO). My data set includes the full set of US utility patents granted between 1836 and 2012, which amounts to 8.3 million patents. The year 1836 marks the beginning of the current patent numbering system, so that my dataset includes patent #1.⁷

Patents represent an imperfect but widespread and voluminous record of innovation outputs. To secure a patent grant, an idea is assessed on three criteria by a patent examiner. Patents must be novel relative to the prior art, nonobvious to someone with ordinary skills in the field, and useful, in the sense that it solves some problem.⁸ These hurdles make patent grants a useful proxy for effective ideas, in the sense used by this paper.

To be sure, patents are not a perfect measure of ideas. Not all ideas are patented, or even patentable. Abstract ideas, for instance, cannot be patented. Moreover, ideas that are *patentable* may not be *patented*, even if they are novel, nonobvious, and useful, because patenting is not costless and

⁷ Prior to the 1836 numbering scheme, an additional 9,957 patents were issued, which are not in my dataset. See US Patent and Trademark Office (2014a).

⁸ See Clancy and Moschini (2013).

requires divulging the details of the innovation.⁹ Furthermore, patents can also be inappropriately granted if the patent office is too lenient.¹⁰ That said, no alternatives are clearly superior and patents provide an extensive source of microdata on individual innovations across a huge array of industries and years. No other single source provides coverage across hundreds of sectors and more than 100 years.

5.2. Technology Mainlines as Elements

With patents proxying for the realization of an effective idea, a decision to be confronted is what to use as a proxy for the elements that are combined to build ideas. I choose to use the technology classifications assigned to each patent. The USPTO has developed the US Patent Classification System (USPCS) to organize patent and other technical documents by common subject matter. Subject matter can be divided into a major component called a class, and a minor component, called a subclass. The USPTO states “A class generally delineates one technology from another. Subclasses delineate processes, structural features, and functional features of the subject matter encompassed within the scope of a class.”¹¹ Subclasses are therefore a natural candidate for the elements of combination, out of which are built new ideas.

Patent subclasses as proxies for elements of combination have many advantages over plausible alternatives, such as the words used in a patent document or citations to prior art. Unlike text or citations, patent classifications are chosen by an ostensibly disinterested party, namely the patent examiner. Classifications have no special legal standing and are not generally of interest to patent applicants (and therefore not chosen strategically). Instead, they are chosen to facilitate searches by future parties who wish to verify that new applications are, in fact, novel. These classifications are also meant to be exhaustive and non-overlapping, two desirable characteristics for our proxy. Finally, the classification system is updated over time, with older patents assigned updated classifications as the system changes, so that searches of the patent record remain feasible. In contrast, the words used to describe common features may change with legal and aesthetic fashion, but are not retroactively updated as the nomenclature changes.

There are more than 450 classes and more than 150,000 subclasses in the USPCS. To take two examples, class 014 corresponds to “bridges,” and class 706 corresponds to “data processing

⁹ See Clancy and Moschini (2013).

¹⁰ See Bessen and Meurer (2008) for a study related to software patents.

¹¹ US Patent and Trademark Office (2012b), pg I-1.

(artificial intelligence).” A complete list of the current classes can be found on the USPTO website.¹² The subclasses are nested within each class, and correspond to more fine-grained technological characteristics. For example, subclass 014/8 corresponds to “bridge; truss; arrangement; cantilever; suspension,” while the subclass 706/29 corresponds to “data processing (artificial intelligence); neural network; structure; architecture; lattice.”

Subclasses are nested and hierarchical. The uppermost subclass is called a mainline subclass, hereafter simply “mainline.” For example, the subclasses “bridge; truss,” and “data processing (artificial intelligence); neural network,” are both mainlines. The subclass nested one level down is said to be “one indent” in from the mainline. For example, the subclass “bridge; truss; arrangement” is one indent in from the mainline “bridge; truss.” Within these one indent subclasses will be still further subclasses, called two indent subclasses, and so on.

Definition 11: Mainline. The USPCS subclass one indent in from a class.

A classification is assigned to a patent by the patent office with the following methodology.¹³ The examiner has some portion of the patent, called the subject matter, he would like to assign to a subclass. Scanning through the list of mainlines in a class, the examiner stops when he finds a mainline that corresponds to the subject matter. The examiner then scans through the list of subclasses one indent in from the mainline. If none of the one indent subclasses apply to the subject matter, the examiner assigns the mainline to the patent. If one of the subclasses does apply, then the examiner repeats this process for the two indent subclasses that lie within the one indent subclass. The examiner then repeats the above process for three indent subclasses and so forth, until he arrives at a point where no deeper subclasses apply to the subject matter. At this point, the highest indent subclass (which will correspond to the most specific and narrow definition) found to be applicable is assigned to the patent. The USPCS is continually updated to reflect new technological categories, and patent classifications are updated as part of this process.

For every patent in my dataset, I observe both the year it was granted and the technology subclasses to which it is assigned. However, simply using the technology subclasses as elements to be combined is problematic because the categories may not correspond to the same level of

¹² <http://www.uspto.gov/web/patents/classification/selectnumwithtitle.htm>

¹³ US Patent and Trademark Office (2012b), page I-13.

specificity, since they are nested. For example, consider three subclasses, that all belong to class 706, “data processing (artificial intelligence).”

- 706/29 – Data processing (artificial intelligence); neural network; structure; architecture; lattice.
- 706/15 – Data processing (artificial intelligence); neural network.
- 706/45 – Data processing (artificial intelligence); knowledge processing system.

Classes 706/29 and 706/15 are both associated with neural networks, but at different levels of specificity, while 706/45 is not associated with neural networks at all. Without looking at the USPC index, we would not know there is a relationship between some of the subclasses, but not others.

Instead, I use technology mainlines as my primary elements of combination. This identifies a set comprising approximately 17,000 elements. Of these, approximately 13,000 are assigned to utility patents in my dataset (from here on, I restrict attention to the set of mainlines actually assigned). These mainlines are meant to be exhaustive and nonoverlapping. The mean number of mainlines used per class is 29.6, with a median of 21.

The number of mainlines per class varies widely. The maximum is 246 mainlines in one class, while 18 classes have just 1 mainline each. If each mainline is assumed to proxy for a specific technological building block, then ideally each mainline would cover roughly the same scope of technologies. However, we must remember that the technology classification system itself differs in how many technologies are encompassed in one class. For instance, the class 002, which corresponds to “apparel,” has 33 mainlines in it. Class 004, which corresponds to the group “baths, closets, sinks, and spittoons,” has 70. Because class 004 appears to tie together a more disparate set of technologies, more mainlines does not necessarily indicate that a mainline from class 002 covers a wider set of technologies than a mainline from class 004. Moreover, in some cases, the reverse happens and one type of technology is split into many classes. For example, classes 532-570 all correspond to organic compounds, and classes 520-528 all correspond to synthetic resins or natural rubbers. Of the 18 classes with one mainline each, 13 belong to one of these two series. In these cases, classes already divide up the space of technologies very finely, so that additional division into many mainlines is not necessary. Because defining the scope of what constitutes an element across different technologies is bound to be somewhat arbitrary, using mainlines as a proxy seems to me an appropriate first step.

5.3. Assigning Each Patent A Combination of Mainlines

The USPTO makes available a large text document (U.S. Patent and Trademark Office 2014c) listing the technology subclass assigned to each patent, under the most recent classification scheme.¹⁴ Each line of this text document contains a patent number, a subclass code, and an indicator for whether the subclass is the primary subclass (discussed more in Section 10). I extract from this document the subclasses assigned to each patent, as well as the identity of the primary subclass. I can then use the patent number to infer the year of the patent's grant.¹⁵ For the reasons discussed above, I next collapse each technology subclass down to the mainline to which it belongs. For example, US patent 7,640,683 is titled "Method and apparatus for satellite positioning of earth-moving equipment" and describes a method of attaching antenna to the arm of an earthmoving machine in such a way that using satellite positioning systems is possible. This patent was assigned to four technological subcategories:

1. 37/348 – Excavating; ditcher; condition-responsive.
2. 414/699 – Material or article handling; vertically swinging load support; shovel or fork type; tilting; control means responsive to sensed condition
3. 701/50 – Data processing: vehicles, navigation and relative location; vehicle control guidance, operation, or indication; construction or agricultural vehicle type
4. 37/382 – Excavating; road grader-type; condition responsive.

Using the USPC index¹⁶ I coded a program to reassign each subclass to its associated mainline. Applying this program to the above patent, I reclassify it as consisting of the following elements:

1. 37/347: Excavating; ditcher
2. 414/680: Material or article handling; vertically swinging load support

¹⁴ Technology classifications can be freely downloaded from <http://patents.reedtech.com/classdata.php>. I downloaded it in August 2014.

¹⁵ This can be inferred from US Patent and Trademark Office (2014a).

¹⁶ Specifically, I download the US Manual of Classification file from <http://patents.reedtech.com/classdata.php>, in August 2014.

3. 701/1: Data processing: vehicles, navigation and relative location
4. 37/381: Excavating; road grader-type

Ideally, every patent would be comprised of two or more mainlines, because the model is premised on ideas as sets of pairs. In practice, after collapsing all patent assignments to mainlines, only 62.5% of patents are assigned more than one mainline. This share varies over time, averaging 40% over the period 1836-1935 and 69% between 1936 and 2012. A potential explanation is that new classification schemes consolidate commonly used pairs of mainlines into a single technology subclass. If this is the case, then we do not observe as many combinations of elements for older technologies, because combinations frequently used are re-classified as a single technology. I hope to explore this potential explanation in future work.

Over the entire period, the mean number of mainlines per patent is 2.29. The share of patents in a given year assigned more than one mainline, as well as the mean number of mainlines per patent, are each plotted in Figure 1. Over time, the number of elements in a patent has grown. While this may be the consequence of consolidation of subclasses, as discussed above, an alternative interpretation is that the complexity of patents has risen over time. The behavior of complexity in the context of this model is the subject of another research project.

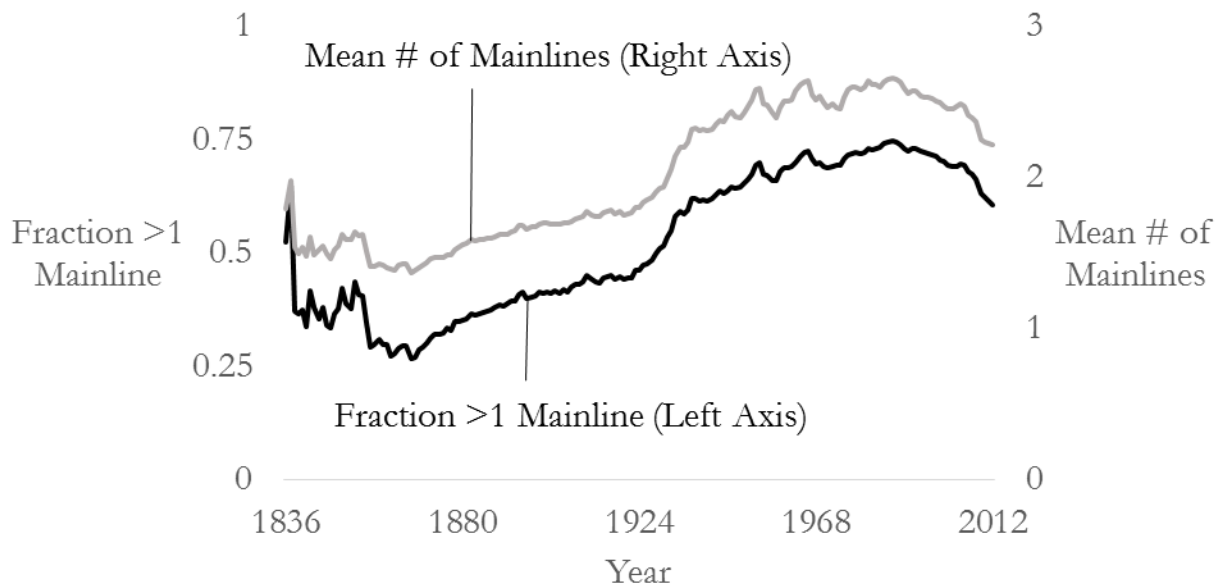


Figure 1: Mainlines per Patent

There is also a converse problem of more than one patent being assigned the exact same set of mainlines. In the model, an idea is defined by the combination of elements that comprise it, so that no two ideas can share the same set of elements without being identical. However, mainlines are aggregations of multiple distinct technological features, since I have collapsed subclasses into mainlines. An exact analogue for an element, in the sense used by the model, is infeasible. However, so long as firms learn from observing combination of mainlines, and so long as the number of distinct technologies drawing on the same mainlines is finite (and small), then our predictions remain apt. I believe these are reasonable assumptions.

Out of approximately $13,000 \cdot 13,000 / 2 = 84.5$ million possible mainline pairs, 1.98 million pairs are actually assigned to at least one patent over the period 1836-2012. Put another way, of the 84.5 million possible pairs, I observe about 2% of the pairs as belonging to effective ideas. The mean number of patents each pair belongs to over the entire period is 10.9, but the distribution is highly skewed. Some 50.6% of observed pairs are only ever assigned to one patent, but the maximum patents assigned to a pair is 22,113.

To summarize, my dataset comprises 8.3 million patents granted between 1836 and 2012. Each of these patents is represented as a combination of mainlines, with 62.5% of patents being assigned more than one mainline. Since patents proxy for the realization of an effective idea, when two mainlines are assigned to the same patent, this proxies for one observation of compatibility between the mainlines.

5.4. Timing Issues

Another mismatch between the model and the dataset involves timing issues. Our data includes the year a patent was granted, but has no information on the year the patent applicant decided to initiate research on the idea that was subsequently patented. I want to test our predictions using only information that would plausibly have been available at the time the researcher decided to begin research, and so will construct the independent variables from lagged data.

The choice of lag length depends on the time required to get a patent granted, the time needed to conduct research, and the speed of diffusion about the outcomes of other researcher projects. Data on patents from 1967-2000 from Hall, Jaffe, and Trajtenberg (2001) suggests the typical lag between a patent application and grant is 2 years. Suppose we assume, optimistically, that research takes 1 year to complete and the results of research are instantly made available to all rival firms. In this case, a patent granted in year t was applied for in year $t - 2$, and the decision to initiate research

began in year $t - 3$. Such a researcher would be able to base his decision on any information available in year $t - 3$. If the results of rival firm projects are instantly diffused, then this includes the outcomes of all research projects initiated before year $t - 4$. Projects initiated in year $t - 4$ are finished in year $t - 3$ and patents are granted in year $t - 1$. So patents granted in year t can draw on information derived from patents granted in year $t - 1$. Hence, we might construct our independent variables from data lagged by one year.

Conversely, suppose, pessimistically that research takes 6 years to complete (the time needed to finish a long Ph.D.) and research is only made public by patent grants themselves. In this case, a patent granted in year t is applied for in year $t - 2$ and research is initiated in year $t - 8$. If the outcome of rival firm research is only revealed by patent grants, then patents granted in year t can draw on information derived from patents granted in year $t - 8$. Hence, we might construct our independent variables from data lagged by eight years.

These arguments provide some plausible bounds for lagged variables, and I experiment with a few variations. In practice, I find lags of 3-5 years provide the best fit.

6. Probability of Use

In Section 4, Predictions 3-5 pertain to the probability a given pair is used, where use means a research project is conducted over an idea that includes the pair. To test these predictions, I draw a sample of data from the patent dataset and use a logit model to test the probability a pair of mainlines is used in any given year. Section 6.1 will present the function form and variables used to assess these predictions, section 6.2 will discuss the results, and section 6.3 will explore the robustness of these results.

6.1. Functional Form and Explanatory Variables

To test prediction 3-5, I draw a sample of 10,000 pairs of mainlines from the 1.84 million that are assigned to at least one patent (section 6.3 will broaden the sample to pairs never used). Using this sample, I estimate the following logit model:

$$\Pr(u_{p,t} = 1) = \text{logit} \left(\sum_i \phi_i (\text{Compatibility Dummy})_{i,p,t-l} + \beta_1 \cdot \text{Age}_{p,t-l} + \beta_2 \cdot \text{Complements}_{p,t-l} + X' \beta \right) \quad (13)$$

where each observation corresponds to a pair-year and the dependent variable $u_{p,t}$ is a dummy variable equal to 1 if the pair of mainlines p is assigned to at least one patent in year t . Notice that all the explanatory variables are lagged by l years. I discuss the explanatory variables below, and their relationship to Prediction 3-5.

Prediction 3 says researchers are more likely to use pairs that have a high expected affinity. Unfortunately, I do not observe the beliefs of researchers, and so I do not know the expected affinity of a given pair. However, expected affinity is increasing in the number of prior observations of compatibility and this can be proxied by the number of patent grants that have been assigned the pair. This is captured by our first set of explanatory variables, which I have called Compatibility Dummy.

The relationship between expected affinity and the number of prior uses of a pair is not necessarily linear though. Indeed, Bayesian updating predicts expected affinity should asymptote at a value of 1, as the number of prior instances of success grow very large. To allow for this, I break prior instances of compatibility into five bins: [1], [2], [3,4], [5,10], and $[11, \infty)$. With the exception of the [1] bin, approximately 10% of pair-year observations fall into each of these bins. I then define a set of dummy variables, (Compatibility Dummy) $_{i,p,t}$, equal to 1 if the number of prior uses of pair p before or up to year t falls into bin i . I predict the coefficients on these bins is increasing and asymptotes for high levels.

Prediction 4 says researchers are less likely to use pairs over time, as they use up all the profitable ideas. To assess this prediction, I define the explanatory variable $\text{Age}_{p,t}$ which counts the number of years that have elapsed between year t and the year in which pair p was first used. I predict the coefficient on $\text{Age}_{p,t}$ will be negative.

Prediction 5 says researchers are more likely to use pairs that are complementary with many other eligible ideas. To assess this prediction, I construct the variable $\text{Complements}_{p,t}$ which proxies for the number of pairs with a comparatively high expected affinity that might belong to eligible ideas. The construction of this variable is more complex than the other two variables.

If I knew the true $E[a(p)]$ of every pair, I could compute $\prod_{p \in d \setminus p'} E[a(p)]$ for every eligible idea, and find some way of aggregating these into a summary statistic. I cannot actually perform such a construction, for practical and computational reasons. To begin, the number of potentially eligible

ideas is literally astronomical (with 13,000 mainlines in use, the number of potential combinations of mainlines is $2^{13,000}$). To get the number of potential ideas down to manageable levels, I restrict my attention to combinations of just 3 mainlines. If each pair can be matched with one of 13,000 other mainlines to form a combination of 3 mainlines, then there are 13,000 such ideas to assess for each pair.

For each pair p , there are approximately 13,000 potential 3-mainline ideas, each of which is defined by one mainline in addition to the pair p . For clarity of exposition, define some idea d which is composed of the mainlines p_1 and p_2 , which jointly form the pair p' , and an additional mainline m . For the idea composed of p_1 , p_2 , and m , $\prod_{p \in d \setminus p'} E[a(p)]$ corresponds to the product of the expected affinities of the pairs corresponding to (p_1, m) and (p_2, m) . I do not observe $E[a(p)]$, but I do observe any patents that have been assigned the mainline pairs (p_1, m) or (p_2, m) . To begin the construction of $\text{Complements}_{p',t}$, I count the number of mainlines m for which there has been at least one previous patent assigned the pair (p_1, m) and a patent assigned the pair (p_2, m) in any year up through year t . As time goes by, this gives an increasing count of the number of mainlines that the pair p' has been connected with from both elements of the pair.

However, simply counting the number of mainlines that have been previously connected with each mainline in a pair will tend to overstate the number of complementary pairs, since eligible ideas are used up over time. My dataset does not allow us to observe ideas that are attempted but which prove ineffective (and are never patented). Instead, I assume that whenever a pair is used, all other eligible ideas that are profitable to attempt are also attempted. I therefore reset the count of complementary pairs to zero each time the pair is used. In section 6.3, I will drop this assumption as a robustness check, and construct a second explanatory variable called $\text{Cumulative Complements}_{p',t}$.

To summarize, the explanatory variable $\text{Complements}_{p',t}$ counts the number of *new* mainlines m that have been used in conjunction with each element in the pair p' since the last time the pair was used. Figure 2 presents the evolution of $\text{Complements}_{p',t}$ for one sample. I predict that the coefficient on $\text{Complements}_{p',t}$ will be positive.

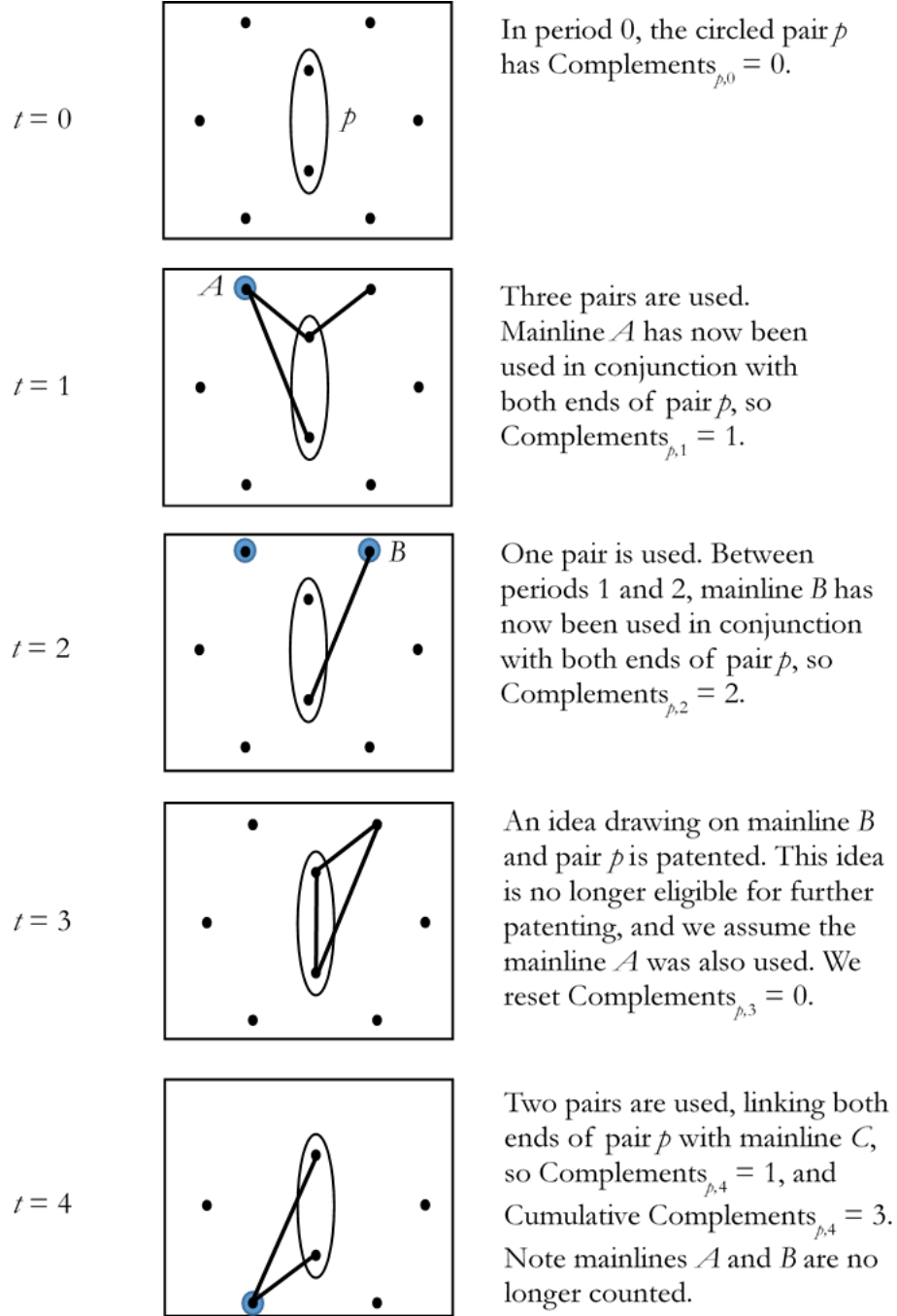


Figure 2: Construction of Complements Variable

Finally X denotes some control variables that I will use in my robustness checks, discussed in section 6.3.

Table 1 presents some summary statistics for the data sample used.

Table 1: Summary of Pair Sample Data

	Min	Median	Mean	Max	St. Dev
$u_{p,t}$	0	0	0.085	1	0.279
Compatibility Dummy					
[1]	0	1	0.556	1	0.497
[2]	0	0	0.143	1	0.350
[3,4]	0	0	0.115	1	0.319
[5,10]	0	0	0.092	1	0.289
[11,∞)	0	0	0.094	1	0.292
Age	0	33	39.67	176	31.22
Complements	0	17	32.42	762	44.02
Cumulative Complements	0	74	111.5	2,113	124.5
$\max\{n_t(p_1), n_t(p_2)\}$	0	43	124.2	6,654	269.5
$\min\{n_t(p_1), n_t(p_2)\}$	0	5	15.89	3,403	42.88
$n_t(p_1) \times n_t(p_2)$	0	184	5,965	16,010,995	77,078

6.2. Results

Table 2 presents the regression results for equation (13).

Column (1) is our baseline result, and the simplest test of Prediction 3-5. Explanatory variables are lagged by 3 years (results are not sensitive to the choice of lag). All the coefficients are statistically significant, and in the directions expected. The probability a pair will be used in any given year is increasing in the number of prior uses, which is consistent with firms learning the expected affinity of the pair is high. The probability of use declines with time, consistent with firms exhausting the best ideas that use any given pair. And pairs are more likely to be used if they have many complements, consistent with this combinatorial model of innovation.

Furthermore, the coefficients on our compatibility dummies strongly support a non-linear relationship between prior uses and probability of use, consistent with a Bayesian learning framework. This is most easily seen in Figure 3, which plots the estimated coefficient for the Compatibility Dummies, against the median number of prior uses for observations that lie within dummy's bin.

This non-linear shape implies researchers learn disproportionately more from their first observations than from later ones.

Table 2: Probability of Use Logistic Regression Results

Dependent Variable	$u_{p,t}$					
	(1)	(2)	(3)	(4)	(5)	(6)
Compatibility Dummy						
[1]	-2.982*** (0.012)	-2.869*** (0.011)	-2.796*** (0.012)	-	-	-
[2]	-2.109*** (0.017)	-2.055*** (0.017)	-1.993*** (0.017)	0.939*** (0.028)	-	-
[3,4]	-1.461*** (0.016)	-1.462*** (0.016)	-1.389*** (0.017)	1.614*** (0.030)	-	-
[5,10]	-0.471*** (0.016)	-0.561*** (0.016)	-0.460*** (0.017)	2.629*** (0.033)	-	-
[11, ∞)	1.579*** (0.018)	1.278*** (0.019)	1.513*** (0.020)	4.410*** (0.043)	-	-
Age	-0.031*** (0.0002)	-0.030*** (0.0002)	-0.029*** (0.0002)	-0.046*** (0.0004)	-	-
Years Feasible	-	-	-	-	-0.0004 (0.002)	-0.009*** (0.002)
Complements	0.007*** (0.0001)	-	0.004*** (0.0001)	0.021*** (0.0002)	0.015*** (0.001)	0.017*** (0.001)
Complements (cumulative)	-	0.001*** (0.00004)	-	-	-	-
Controls	N	N	Y	Y	N	Y
Fixed Effects	N	N	N	Y	N	N
Sample – Pair Use	1 $\geq$	1 $\geq$	1 $\geq$	1 $\geq$	0 $\geq$	0 $\geq$
Observations	564,531	564,531	462,093	462,093	773,233	328,428
Log Likelihood	-126,122	-125,439	-118,338	-88,826	-1,580	-1442

*** signifies significance at $p = 1\%$. Standard errors in parentheses.

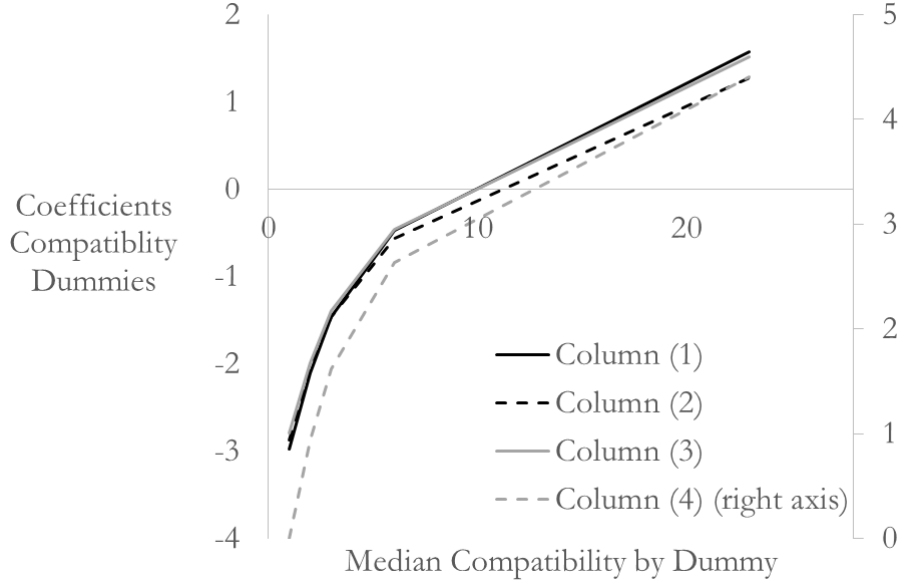


Figure 3: Coefficients of Compatibility Dummy

6.3. Robustness

Columns 2-6 report various robustness checks. In all cases, the results support Prediction 3-5, where applicable. For clarity, I will continue to let p_1 and p_2 denote the mainlines that make up the pair p' .

Column 2 tests a simplified version of the explanatory variable $\text{Complements}_{p',t}$. In the baseline, I reset the count of mainlines that have been used with p_1 and p_2 every time the pair p' is used. This resetting is meant to control for the gradual using up of eligible ideas of which p' form a part. This is a strong assumption, and to check the robustness of the results to its weakening, I construct the explanatory variable $\text{Cumulative Complements}_{p',t}$ which does not perform this resetting. Instead, this variable is a count of all mainlines that have been used with p_1 and p_2 up through period t . Its cumulative nature means it can never decrease. Using this explanatory variable instead of the baseline's has little overall effect on our results. Note the coefficient on $\text{Cumulative Complements}_{p',t}$ is below the coefficient on $\text{Complements}_{p',t}$, though, which suggests the models fit decreases if I do not attempt to control for the gradual exhausting of eligible ideas.

Columns 3 and 4 attempt to address omitted variable bias by inclusion of additional controls. Let $n_t(p_i)$ denote the number of patents granted in year t that include the mainline p_i . In columns 3 and 4, I add to equation (13) the variables $\max\{n_t(p_1), n_t(p_2)\}$, $\min\{n_t(p_1), n_t(p_2)\}$ and $n_t(p_1) \times n_t(p_2)$. These are meant to control for any time-varying propensities to develop patents that use mainlines p_1 or p_2 outside of this model, and to therefore use the two mainlines together simply by chance. Note that when $n_t(p_i) = 0$, then there are no patents that are assigned p_i and so $n_t(p_1) \times n_t(p_2)$ and $u_{p,t}$ must each be equal to zero. Since I am only interested in predicting probabilities when use is in principle possible, I omit these observations, which reduces the number of observations from 564,531 to 462,093. Introducing these controls reduces the size of the estimated coefficients, but results remain significant and in the directions predicted.

Column 4 adds to this specification fixed effects, which should control for time-invariant propensities to use a given pair. In a non-linear setting, computing fixed effects by de-meaning the data is not appropriate. Moreover, simply estimating a coefficient for dummy variables associated with each pair leads to the incidental parameters problem, which can bias coefficients.¹⁷ Fortunately, the Chamberlain estimator is an alternative to demeaning data that can be used in a logistic regression framework. For every pair p , the Chamberlain estimator conditions estimation on the sum of $u_{p,t}$ over the pair's lifecycle, and it can be shown that this approach strips out fixed effects, just as de-meaning the data does in a linear model.¹⁸ Adding these fixed effects strengthens our results.

Columns 5 and 6 consider sample selection issues. The preceding columns rely on a dataset that is drawn from pairs that are used at least once. This should not bias our results, because I only include observations that begin *after* the initial use of the pair (beginning after the first use allows us to construct all our explanatory variables). However, this sampling methodology implies our results pertain only to the probability of use, after a first initial use. In columns 5 and 6, I see if our model is also applicable to the first time a pair is used.

To assess this, I draw 8,003 pairs of mainlines from the set of 84.5 million possible pairs. Most pairs are never used. From the 8,003 pairs selected, only 183 pairs are ever used by 2012 (2.2%). Since I have an observation for every year and pair, there are 773,233 total observations. I am only

¹⁷ See Greene (2008), pg 800-801.

¹⁸ See Greene (2008), pgs. 803-805.

trying to predict the first time a pair is used now, and so the dummy variables associated with prior uses are all equal to zero, and $\text{Age}_{p,t}$ is undefined. Instead, I find the first year that both mainlines have been assigned to a patent, and set this year as the pair’s first “available” year. I interpret this year as an indicator of the earliest time the pair could possibly have been used by innovators. For example, if the mainline 14/3 (“Bridge; Truss”) was first assigned to a patent in 1836 and the mainline 706/15 (“Data processing (artificial intelligence); neural network”) was first assigned to a patent in 1980, then the pair (14/3, 706/15) is first available in 1980, rather than 1836. I then measure the years the patent has been available. It is worth including this metric, since not using a pair shortly after it is available is a signal that the ideas using the pair do not have $\pi(d)$ high enough or $k(d)$ low enough to warrant a research project even when expected efficacy of the ideas is low.

Without information on prior uses, only Prediction 5 is relevant. For each pair-year, I compute the variable $\text{Complements}_{p',t}$, as before. In column 5, this variable is a strong predictor of first use, which is again consistent with theory. In column 6, I include as additional controls the same ones constructed for column 3. In this instance the results are stronger when these controls are included.

These robustness checks provide further support for the veracity of predictions 3-5. These predictions pertained to the use of specific pieces of knowledge, in this case pairs of mainlines. In the next section, I show that this model’s predictions also apply to entire technology sectors.

7. Patents Issued Per Year

Section 4 argued that more patents will be granted when the expected affinity of pairs in a sector are high, but that the number of patents granted will decrease over time. There are a number of data challenges that must be surmounted in order to test these predictions. I discuss these data issues, as well as some functional form assumptions, in Section 7.1. In Section 7.2, I present my baseline results. Section 7.3 provides some robustness checks.

7.1. Data and Functional Form Issues

7.1.1. Units of Observation

Predictions 1 and 2 pertain to the number of patents granted within a given technology sector, which was understood to be a group of firms drawing upon the same set of technological building blocks. In this section, I use 429 technology classes in the USPCS as proxies for different

technological sectors. This gives us a panel of 429 classes, each with up to 176 years of innovative activity.

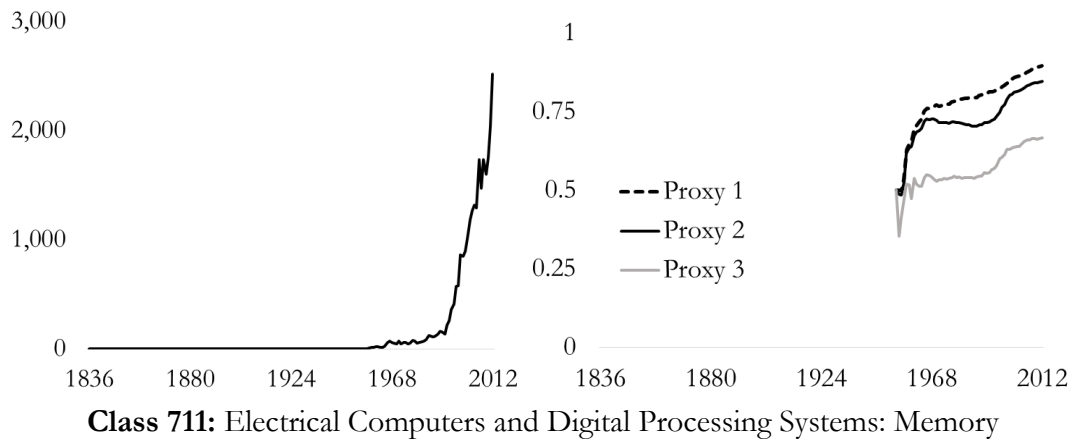
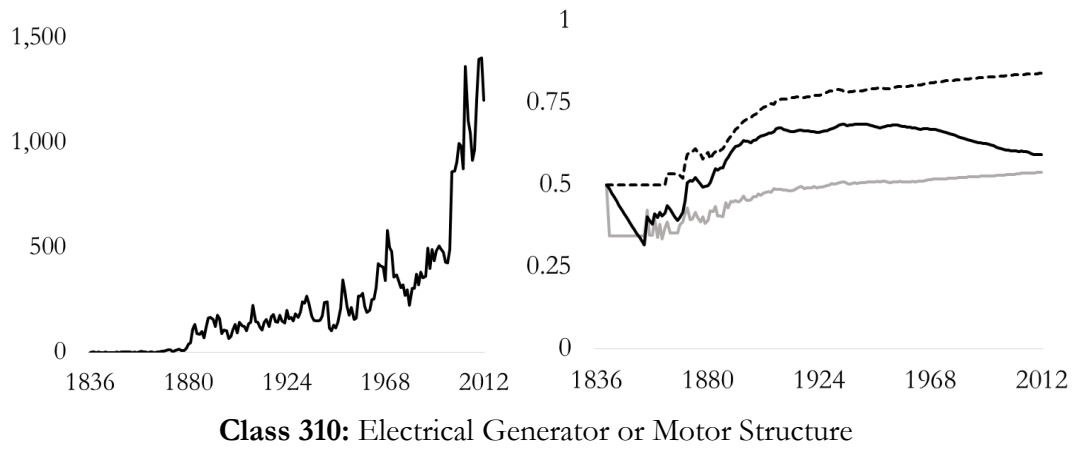
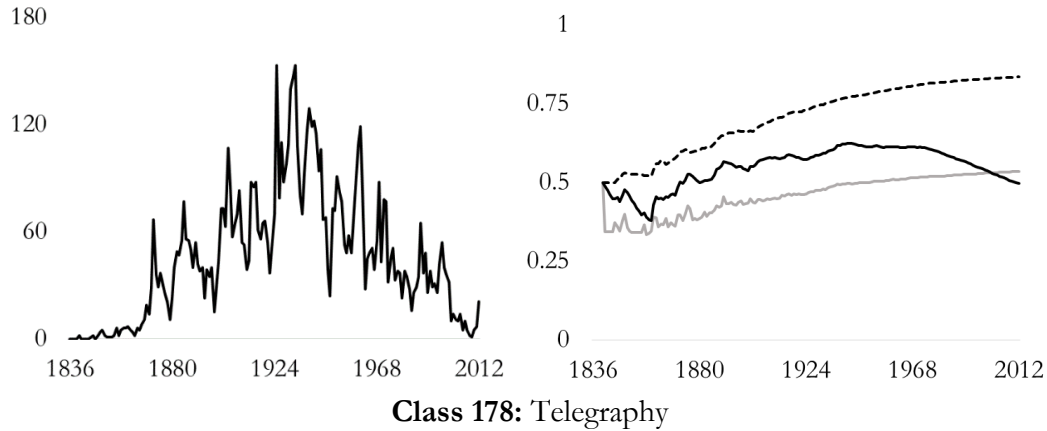
To determine the number of patents granted to a class in a year, I use the primary classification assigned to each patent. Each patent is assigned one, and only one, primary classification, which is based on the main inventive concept. The primary classification is generally used in economics to assign patents to different technology classes (see for example, Hall, Jaffe, and Trajtenberg 2001). The dependent variable $n_{cl,t}$ corresponds to the number of patents with primary class cl granted in year t . The path of three $n_{cl,t}$ variables over time is plotted in the left column of Figure 4.

Testing predictions 1 and 2 also requires information on the pairs used by each technology class. To assemble a list of the pairs used by each class, I tally the mainlines assigned to each patent in a class. For example, recall that patent 7,640,683, for “Method and apparatus for satellite positioning of earth-moving equipment” was assigned the mainlines:

5. 37/347: Excavating; ditcher
6. 414/680: Material or article handling; vertically swinging load support
7. 701/1: Data processing; vehicles, navigation and relative location
8. 37/381: Excavating; road grader-type

The primary class for this patent is class 37, excavating, and it was granted in year 2010. I therefore assign mainline-pairs between 37/347, 414/680, 701/1, and 37/381 as belonging to class 37 from 2010 onwards. After doing this for every patent in the class, I obtain the list of all mainline-pairs used by any patent in the class, in each year.

The number of mainlines used by a class is different from the number of mainlines that are *nested* under a class by the USPCS, because most patents assigned to a class draw on mainlines from outside the class (just as the above patent is assigned to class 37 but is also assigned mainlines from classes 414 and 701). The mean number of mainlines nested under one class is 29.6, but the mean number of mainlines used by patents in a class is 1,161. Moreover, the minimum number of mainlines nested under one class was 1, while the minimum used by patents in a class is 5. The maximum number of mainlines nested under a class was 246, while the maximum used by a class is 5,126.

Left Column: Patents Granted per Year**Right Column: Average Affinity Proxies****Figure 4: Patent Class Trends**

7.1.2. Proxying for Expected Affinity

Predictions 1 and 3 both relate to the expected affinity of a pair. In section 6, I did not attempt to construct a proxy for expected affinity, but instead took a relatively non-parametric approach to the question. The results were consistent with a Bayesian learning framework, where more observations of success are correlated with a higher probability of use, and the first few such observations are more important than later ones. In this section, I will be a bit more ambitious, and explicitly construct three different proxies for the expected affinity of a pair. Given that we do not observe the beliefs of researchers, nor research projects that do not yield patents, these proxies are bound to be noisy. Nevertheless, as we will see, they do carry useful information.

Each of these proxies is 0 when there are no observations of prior success, and therefore assumes the initial prior of research firms assigns a very low probability any given pair of elements will be compatible. Each of these proxies then draws on past observations of $s_t(p)$, where $s_t(p)$ is the number of patents pair p is assigned to in year t .

The first proxy for expected affinity is:

$$\tilde{E}_t^1[a(p)] \equiv \frac{\sum_{\tau=0}^t s_{\tau}(p)}{1 + \sum_{\tau=0}^t s_{\tau}(p)} \quad (14)$$

This is the simplest proxy, and it exploits only information on the cumulative number of past observations of a pair being assigned to a patent. It approaches 1 as the total number of prior uses of a pair goes to infinity.

The second proxy for expected affinity is:

$$\tilde{E}_t^2[a(p)] \equiv \frac{\sum_{\tau=0}^t 0.95^{t-\tau} s_{\tau}(p)}{1 + \sum_{\tau=0}^t 0.95^{t-\tau} s_{\tau}(p)} \quad (15)$$

This formulation is equivalent to the first proxy, but where more recent observations of a pair are accorded more weight. It implicitly assumes mainlines are an imperfect proxy for the true building blocks that make up an idea, and assumes more recent observations may be more relevant

for estimating expected affinity. In practice, this proxy builds in a tendency for $E[a(p)]$ to decay over time if a pair is not continuously assigned to new patents.

The third proxy for expected affinity is:

$$\tilde{E}_t^3[a(p)] \equiv \frac{\sum_{\tau=0}^t s_{\tau}(p)}{1 + \sum_{\tau=0}^t n_{\tau}(p)} \quad (16)$$

where

$$n_t(p) = \max\{s_t(p), (1 + g_t)s_{t-1}(p)\} \quad (17)$$

and

$$1 + g_t = \sum_p s_t(p) / \sum_p s_{t-1}(p) \quad (18)$$

This formulation attempts to mechanically infer the number of times researchers tried to use a pair with a simple rule. In each period, researchers expect the number of attempts to use pair p is equal to the number of times the pair was assigned to patents in the last period, multiplied by an aggregate growth term $1 + g_t$, or the number of patents it was actually assigned to this period, whichever is larger. The growth term $1 + g_t$ is the overall growth rate of patent pair assignments. In practice, this proxy penalizes pairs that fail to “keep up” with the aggregate growth rate of all pair assignments, by assuming their failure to keep up reflects a failure of research ideas to be effective, rather than an absence of research attempts.

As a simple test of these three proxies, I use them in place of the Compatibility Dummies on the same sample of pairs drawn in Section 6 to estimate the following:

$$\Pr(u_{p,t} = 1) = \text{logit}\left(\beta_0 \cdot \tilde{E}_{t-l}^i[a(p)] + \beta_1 \cdot \text{Age}_{p,t-l} + \beta_2 \cdot \text{Complements}_{p,t-l} + \beta_3\right) \quad (19)$$

The results to regression (19) are presented in Table 3.

Note that the log-likelihood using the first proxy is no better than the non-parametric baseline, which is plausible, since both embody a concave function of cumulative past observations. The log-likelihood improves significantly when I begin to weight more recent observations more than distant

ones, as in the second proxy, and improves again if I exploit the growth rate of pair uses, relative to all pair uses, as in the third proxy.

Table 3: Probability of Use Logistic Regression Results, With Expected Affinity Proxies

Dependent Variable	$u_{p,t}$		
	Proxy 1	Proxy 2	Proxy 3
$\tilde{E}_{t-l}^i[a(p)]$	9.443*** (0.041)	10.240*** (0.041)	24.982*** (0.094)
$\text{Age}_{p,t-l}$	-0.027*** (0.0002)	-0.020*** (0.0002)	-0.030*** (0.0002)
$\text{Complements}_{p,t-l}$	0.006*** (0.0001)	0.019*** (0.0002)	0.009*** (0.0002)
Observations	564,531	564,531	564,531
Log-Likelihood	-127,581	-95,550	-79,152

*** signifies significance at $p = 1\%$. Standard errors in parentheses.

Given my proxies for $E[a(p)]$, and the assignment of mainline-pairs to classes, I can now compute a measure for the average expected affinity for each class, in each year:

$$(\text{Average Affinity Proxy})_{cl,t-l} = \frac{1}{N_{cl,t}} \sum_{p \in cl,t} \tilde{E}_t^i[a(p)] \quad (20)$$

I plot the average affinity proxy in the right column of Figure 4, for three example classes. As can be seen, my measure for average affinity proxy is only defined from the point at which the first patent in the class is granted. Moreover, it exhibits variation over time. Finally, although the three proxies do not move in sync with each other, they are correlated.

7.1.3. Functional Form

To test predictions 1 and 2, a simple but naïve approach would be to estimate the following regression:

$$n_{cl,t} = \tilde{\beta}_0 + \tilde{\beta}_1 \cdot (\text{Average Affinity Proxy})_{cl,t-l} + \tilde{\beta}_2 \cdot \text{Age}_{p,t-l} + \varepsilon_{cl,t} \quad (21)$$

Where $\text{Age}_{cl,t}$ denotes the number of years that have elapsed since a patent was first given class cl as its primary assignment. The unit of observation in equation (21) is a class-year, and I would

expect the coefficient on the Average Affinity Proxy to be positive and the coefficient on Age to be negative.

There are a number of problems with this naïve specification. First, classes may vary widely in the number of firms, the inherent profitability of research in the sector, or the number of building blocks available. Indeed, as Figure 4 makes clear, the number of patents granted across different classes is highly variable. These differences suggest a log model is more appropriate to estimate. To estimate a log model, I drop the 6.3% of observations where $n_{cl,t} = 0$, so that the results should be viewed as being conditional on positive research activity taking place. In my robustness checks, I add 1 to the observations so that these dropped observations can be included.

Second, it is clear from Figure 4 that the number of patents granted per year and Average Affinity Proxy are trending variables, so that a simple regression of one onto the other would mostly pick up trends. Since these trends may well vary by class, I estimate a fixed effect model on differenced data:

$$\Delta \ln n_{cl,t} = \alpha_{cl} + \beta_1 \cdot \Delta \ln(\text{Average Affinity Proxy})_{cl,t-l} + \beta_2 \cdot \text{Age}_{cl,t-l} + X'_{cl,t} \beta + \varepsilon_{cl,t} \quad (22)$$

Note the term α_{cl} can account for class-specific growth rates.

Third, to account for higher order trends and omitted variables, I include as controls several lags of the dependent variable $\Delta \ln n_{cl,t}$. Moreover, to control for omitted variables that influence aggregate patenting, I include lags of changes in log-aggregate patenting $\Delta \ln N_t$ where $N_t \equiv \sum_{cl} n_{cl,t}$.

To determine the number of lags to include, I estimate the following baseline model, which does not include the variable $\Delta \ln(\text{Average Affinity Proxy})_{cl,t-l}$:

$$\Delta \ln n_{cl,t} = \alpha_{cl} + \sum_{i=1}^8 \varphi_i L^i \Delta \ln n_{cl,t} + \sum_{i=1}^3 \tilde{\varphi}_i L^i \Delta \ln N_t + \beta_3 \cdot \text{Age}_{cl,t} + \varepsilon_{cl,t} \quad (23)$$

This model forecasts changes in class patenting activity with a class-specific fixed effect, lags of changes in class patents, lags of the change in total patents, and the age of the class. The estimated coefficients for this model are presented in Table 4.

Table 4: Coefficients of Baseline Model

Dependent Variable:	$\Delta \ln n_{cl,t}$								
Lags	0	1	2	3	4	5	6	7	8
$\Delta \ln n_{cl,t}$		-0.421 (0.009)	-0.246 (0.008)	-0.178 (0.008)	-0.120 (0.007)	-0.099 (0.005)	-0.066 (0.006)	-0.033 (0.006)	-0.021 (0.005)
$\Delta \ln N_t$		0.446 (0.016)	0.248 (0.016)	0.130 (0.016)					
$Age_{cl,t}$	-0.001 (0.000 04)								
Class Fixed Effects?	Yes								
Observations	53,930								
Adjusted R²	0.145								

Note: All coefficients are significant at $p = 0.1\%$. White standard errors clustered by class in parentheses.

I chose to include 8 lags of $\Delta \ln n_{cl,t}$ and 3 lags of $\Delta \ln N_t$ since the statistical significance levels of further lags is lower, and also reduces the adjusted R². The baseline model has a few notable features. First, there is a tendency for the growth rate of patents in a class to converge to the aggregate growth rate of patents, which can be seen from the approximately equal but opposite coefficients on $\Delta \ln n_{cl,t}$ and $\Delta \ln N_t$. Second, the growth rate of patents slows over time, as indicated by the negative coefficient on $Age_{cl,t}$. Finally, a Hausman test strongly rejects the hypothesis that class fixed effects can be ignored. Different classes have different growth rates. Unless explicitly stated, all the explanatory variables in equation (23) and Table 4 are included in all the following regressions, although I do not report their estimated coefficients to save space.

Lastly, the Average Affinity Proxy conflates two sources of variation, namely the sample of pairs and $\tilde{E}_t^i[a(p)]$ for each pair. In period t , the set of pairs used by a class cl is defined as the set of pairs used by all patents awarded to class cl in any period prior to or including period t . However, defining classes in this manner means the set of pairs used is growing over time, rather than constant. To make sure the change in expected affinity is due to changes in affinity, rather than the definition of the set, in each period, I compute $\Delta \ln(\text{Average Affinity Proxy})_{cl,t-l}$ using the same set of pairs. Specifically, when comparing Average Affinity Proxy in period t and period $t-1$, I always use the set of pairs defined for period t .

7.1.4. Data Description

After computing the above metrics, I have an unbalanced panel of 429 classes with up to 176 years of data per class (the panel is unbalanced, because I exclude years before a class has its first granted patent), for a total of 59,592 observations. Some summary statistics are presented below in Table 5.

Table 5: Summary Patent Class Statistics

	Min	Median	Mean	Max	Standard Deviation
$Age_{cl,t}$	1	81	82.8	176	47.39
Patents Granted ($n_{cl,t}$)	0	57	132.9	8,348	276.1
Average Affinity Proxy 1	0.071	0.722	0.705	0.999	0.092
Average Affinity Proxy 2	0.064	0.543	0.552	0.996	0.090
Average Affinity Proxy 3	0.053	0.455	0.454	0.805	0.060
$\Delta \ln n_{cl,t}$	-2.984	0.001	0.031	4.414	0.426
$\Delta \ln(\text{Average Affinity Proxy 1})_{cl,t-l}$	0.000	0.012	0.034	2.881	0.090
$\Delta \ln(\text{Average Affinity Proxy 2})_{cl,t-l}$	-0.044	0.008	0.030	2.877	0.097
$\Delta \ln(\text{Average Affinity Proxy 3})_{cl,t-l}$	-1.216	0.010	0.030	3.105	0.110

Note that the mean values for the change in log-transformed variables are all of a similar magnitude, although $\Delta \ln n_{cl,t}$ has more variance than any of the $\Delta \ln(\text{Average Affinity Proxy})_{cl,t-l}$ metrics. Some correlations are presented in Table 6:

Table 6: Correlations Among Proxies

	Average Affinity Proxy		
	1	2	3
Average Affinity Proxy 1	1		
Average Affinity Proxy 3	0.919	1	
Average Affinity Proxy 2	0.661	0.786	1

	$\Delta \ln(\text{Average Affinity Proxy})_{cl,t-l}$		
	1	2	3
$\Delta \ln(\text{Average Affinity Proxy 1})_{cl,t-l}$	1		
$\Delta \ln(\text{Average Affinity Proxy 2})_{cl,t-l}$	0.919	1	
$\Delta \ln(\text{Average Affinity Proxy 3})_{cl,t-l}$	0.994	0.931	1

Note that these three measures are much more correlated in their differences than in their levels. This stems from the fact that the difference between Average Affinity Proxy in one period to another is most prominently driven by the pairs which go from 0 to 1 assigned patents, which all three proxies measure as a change from 0 to 0.5.

7.2. Baseline Results

In Section 6, my preferred specification was a lag of 3 years, but I found the choice of lag had little impact on the estimated results. In this model, the choice of lag is not inconsequential. My first investigation examines all three proxies for $\Delta \ln(\text{Average Affinity Proxy})_{cl,t-l}$ across two lag specifications. The results are displayed in Table 7.

As discussed earlier, the significance of explanatory variables with a lag greater than two is consistent with patent grants as a primary vehicle for knowledge diffusion, while the significance of variables with smaller lags indicates knowledge can spread before the patent is granted. I find evidence for both effects, with $\Delta \ln(\text{Average Affinity Proxy})_{cl,t-l}$ usually positive and statistically significant for lags of 1 year and 3-5 years, although significance is stronger for the 3-5 year period. Since research may take different lengths of time for different projects, it is not surprising that multiple years of explanatory variable are statistically significant.

The coefficients on all three proxies are similar in magnitudes, but as found in Table 7, Proxy 3 has the strongest predictive power. It is the only proxy with statistical significance for a lag of 5 years, and the adjusted R-squared is usually highest for this proxy. I also conduct an F-test for the joint hypothesis that all coefficients on lagged values of $\Delta \ln(\text{Average Affinity Proxy})_{cl,t-l}$ are statistically indistinguishable from zero, but this is rejected at $p = 0.1\%$ in every specification. Again, the F-statistic is largest for Proxy 3.

Table 7 supports Prediction 1 and 2. The average expected affinity between elements in a technology field is a positive predictor of the change in patents granted in the future, with an elasticity between 0.1 and 0.2, but applying over several periods. This occurs even though I control for lagged behavior, and despite the imprecision in the proxies for elements, ideas, and expected affinity. Moreover, there is a general tendency for fewer patents to be granted as time goes on.

Table 7: Different Lag and Proxy Specifications

Dependent Variable:	$\Delta \ln n_{cl,t}$					
	(1)	(2)	(3)	(4)	(5)	(6)
	Proxy 1	Proxy 2	Proxy 3	Proxy 1	Proxy 2	Proxy 3
Age	-0.001*** (0.0001)	-0.001*** (0.0001)	-0.001*** (0.0001)	-0.001*** (0.00005)	-0.001*** (0.00005)	-0.001*** (0.00005)
$L\Delta \ln(\text{Average Affinity Proxy})$	0.167** (0.077)	0.120* (0.064)	0.127** (0.053)			
$L^2\Delta \ln(\text{Average Affinity Proxy})$	-0.058 (0.079)	-0.054 (0.066)	0.016 (0.053)			
$L^3\Delta \ln(\text{Average Affinity Proxy})$	0.191*** (0.065)	0.152*** (0.054)	0.142*** (0.043)	0.202*** (0.061)	0.156*** (0.052)	0.149*** (0.043)
$L^4\Delta \ln(\text{Average Affinity Proxy})$	0.173*** (0.059)	0.152*** (0.050)	0.167*** (0.041)	0.186*** (0.059)	0.158*** (0.050)	0.179*** (0.041)
$L^5\Delta \ln(\text{Average Affinity Proxy})$	0.082 (0.060)	0.077 (0.051)	0.116*** (0.041)	0.088 (0.056)	0.078 (0.049)	0.122*** (0.040)
$L^6\Delta \ln(\text{Average Affinity Proxy})$	-0.040 (0.049)	-0.036 (0.042)	0.011 (0.034)			
$L^7\Delta \ln(\text{Average Affinity Proxy})$	0.040 (0.050)	0.029 (0.043)	0.015 (0.034)			
$L^8\Delta \ln(\text{Average Affinity Proxy})$	0.002 (0.044)	0.004 (0.038)	-0.022 (0.032)			
Controls						
Lags of $\Delta \ln n_{cl,t}$	8	8	8	8	8	8
Lags of $\Delta \ln N_t$	3	3	3	3	3	3
Class Fixed Effects	Y	Y	Y	Y	Y	Y
F-Test†	4.925	4.000	5.346	11.086	9.340	11.952
Observations	53,930	53,930	53,930	53,930	53,930	53,930
Adjusted R ²	0.147	0.146	0.147	0.146	0.146	0.147

Note: † The null hypothesis is the joint insignificance of all coefficients on lags of $\Delta \ln \text{Average.affinity.proxy}$.
White standard errors clustered by class are in parentheses. * $p < 0.1$; ** $p < 0.05$; *** $p < 0.01$

7.3. Robustness Checks

I next perform a set of robustness checks. A more detailed discussion is available in the appendix. A primary result is the variable strength of the above-noted relationship over time. Generally speaking, the coefficients on $\Delta \ln(\text{Average Affinity Proxy})_{cl,t-l}$ are most likely to be positive and statistically significant for the period of 1936-2012. In the robustness checks, I will sometimes find the coefficients on $\Delta \ln(\text{Average Affinity Proxy})_{cl,t-l}$ lose significance or change sign, but this result is invariably driven by the 1836-1936 period. The 1936-2012 period generally conforms to the patterns described in section 7.2, and is often stronger.

There are two reasons why this result is not surprising. As shown in Figure 1, the fraction of patents with more than one mainline varies substantially over the period 1836-2012. Since the model is premised on ideas consisting of combinations of at least 2 elements, the model is a better fit for the data after 1936, when 69% of patents were assigned 2 or more mainlines, as opposed to 40% before 1936. This suggests the pre-1936 data may be subject to greater measurement error, which would lead to attenuation bias.

In addition to measurement error, a second potential source of bias may stem from the USPTO's ongoing updating of its classification system. If commonly used pairs of mainlines are eventually consolidated into single mainlines, then the only pairs we will observe in earlier periods will be pairs of mainlines that were *not* subsequently combined many times. This would tend to bias the results, with the bias more severe for earlier periods. Indeed, if I break the period 1936-2012 into two more periods from 1936-1986 and 1986-2012, the results for the earlier period remain very strong, but the size of estimated coefficients for 1986-2012 grows substantially (the coefficient on $L^3 \Delta \ln(\text{Average Affinity Proxy})_{cl,t-l}$ is over 2).

I find a similar effect when I measure $\Delta \ln(\text{Average Affinity Proxy})_{cl,t-l}$ with reference to the pairs used in the previous period, rather than the current period. That is, when comparing the Average Affinity Proxy in period t and period $t-1$, I now use the set of pairs defined in $t-1$. This has the effect of omitting all pairs that were used for the first time in period t , and restricts attention to pairs that have already been used at least once. As reported in the appendix, this leads to a *negative* and statistically significant coefficient for the 1836-1935 data, but generally familiar (and strong) positive and significant coefficients for the 1936-2012 period.

Another potential source of bias in my estimates is the construction of proxies. Since these rely on cumulative counts of patent grants, they will be correlated with the growth rate of patents over time, which is the dependent variable. While I have attempted to control for this by adding 8 lags of $\Delta \ln n_{cl,t}$, it may be that the proxy is picking up the effect of lags beyond 8. To check for this, I double the number of lags. This leads to a strengthening of the results, compared to those found in Table 7, column (4)-(6).

In another specification, I restore observations where $n_{cl,t} = 0$ by transforming the dependent variable to $\Delta \ln(1 + n_{cl,t})$. All variables remain positive and significant in this specification. Lastly, when I omit all classes that contain less than two mainlines, results are largely unaffected.

8. Discussion

The preceding empirical analysis shows that accounting for the combinations of mainlines used in a class of patents does help predict the growth rate of patents by class, as compared to a model composed of lags, fixed effects, and class age. Patenting increases are correlated with new combinations made before the patent was granted. Where the data better fits the model, results are stronger, and modifications to the proxy and data used do not dramatically impact the results. The implied elasticity is on the order of 0.1-0.2 for several years, although this may be attenuated by measurement error. If I rely only on data from 1936 onwards, the elasticity rises.

According to the model presented in this paper, a rise in expected affinity for some pairs in a class increases the expected number of patent grants, because ideas using these pairs are more likely to be effective. The rise in class patent output should come disproportionately from patents that include pairs whose increase led to an increase in Average Affinity Proxy. While I have not tested for this effect directly, this story is consistent with the results from section 6, which showed the probability a pair of mainlines is assigned to a patent is also increasing in the expected affinity of the pair, and of the pairs in related ideas.

Across all specifications, I also find the age of a class has negative and significant impact on the growth rate of patents granted. This is also consistent with the presence of fishing out effects, implying a long-run decline in the growth rate of any one class. As a back of the envelope calculation, the growth rate of 2.7% for all patents over the last 100 years falls to zero in 27 years all else equal, when the coefficient on age is -0.001 (a common result). I find a similar time-scale at work when estimating the probability a pair will be used in any given year. In both cases, I find a

story consistent with the model, where technological progress in any given class must constantly reinvigorate itself by discovering new pairs have a high affinity, or else research productivity falls to zero.

9. Conclusions

This paper has presented a new model of knowledge production, and shown some of its predictions are consistent with tests using U.S. patent data. This approach helps clarify some questions about the production of knowledge.

Have we run out of things to invent, or are we on the cusp of a new era of technological advance? This paper argues these two forces are in a constant state of ebb and flow. In the short and medium run, firms fish out the best ideas that can be constructed from a given set of technological building blocks. At the same time, this process helps firms learn which types of combinations are feasible, which serves to expand the set of ideas that are worth their research costs. The future for technological innovation is brightest after a set of technological building blocks has been discovered (possibly via recombinant growth, as in Weitzman 1998, or via other processes) *and* some initial work has established many elements in this set have a high affinity, but *before* the best ideas have been fished out. If a successful sector keeps recycling the same combinations however, it will learn less and less useful information from each invention, and eventually exhaust the finite set of possible discoveries.

The empirical methodology used in this paper could be used as a method of establishing a measure for the outlook for many technological sectors, or academic disciplines. This could then be used as a way to control for technological opportunity in many other applications, or even as a method of forecasting. Such an approach might give some guidance on the question of whether US and global innovation is stagnating or surging.

The approach used in this paper could also serve as one input into funding decisions for agencies conducting R&D across many disciplines and sectors. This paper's model also suggests public funders of science can have the most positive spillovers by doing the exploratory work of bringing together untried combinations of elements.

References

- Arthur, W. Brian. 2009. *The Nature of Technology: What It Is and How It Evolves*. New York: Free Press.
- Auerswald, Philip, Stuart Kauffman, Jose Lobo, and Karl Shell. 2000. "The Production Recipes Approach to Modeling Technological Innovation: An Application to Learning By Doing." *Journal of Economic Dynamics and Control* 24 389-450.
- Bessen, James, and Michael J. Meurer. 2008. *Patent Failure: How Judges, Bureaucrats and Lawyers Put Innovators at Risk*. Princeton, NJ: Princeton University Press.
- Borges, Jorge Luis. 1962. *Ficciones* (English Translation by Anthony Kerrigan). New York, NY: Grove Press.
- Brynjolfsson, Erik and Andrew McAfee. 2014. *The Second Machine Age: Work, Progress, and Prosperity in a Time of Brilliant Technologies*. New York, NY: W.W. Norton & Company, Inc.
- Clancy, Matthew. 2015. "Optimal Research Strategies When Innovation is Combinatorial." *Dissertation*.
- Clancy, Matthew and GianCarlo Moschini. 2013. "Incentives for Innovation: Patents, Prizes, and Research Contracts." *Applied Economic Perspectives and Policy* 35 (2) 206-241.
- Cowan, Tyler. 2011. *The Great Stagnation: How America Ate All The Low-Hanging Fruit of Modern History, Got Sick, and Will (Eventually) Feel Better*. New York, NY: Penguin Group.
- Dartnell, Lewis. 2014. *The Knowledge: How to Rebuild Our World From Scratch*. New York, NY: Penguin Press.
- Fleming, Lee. 2001. "Recombinant Uncertainty in Technological Search." *Management Science* 47(1) 117-132.
- Gordon, Robert J. 2012. "Is U.S. Economic Growth Over? Faltering Innovation Confronts the Six Headwinds." *NBER Working Paper* 18315.
- Greene, William H. 2008. *Econometric Analysis: Sixth Edition*. Upper Saddle River, New Jersey: Pearson Prentice Hall.
- Hall, Bronwyn H., Adam B. Jaffe, and Manuel Trajtenberg. 2001. "The NBER Patent Citation Data File: Lessons, Insights and Methodological Tools." *NBER Working Paper* 8498.
- Harhoff, Dietmar, Francis Narin, F.M. Scherer and Katrin Vopel. "Citation Frequency and the Value of Patented Inventions." *The Review of Economics and Statistics*, 1999: 511-515.
- Jovanovic, Boyan and Rafael Rob. 1990. "Long Waves and Short Waves: Growth Through Intensive and Extensive Search." *Econometrica* 1391-1409.
- Kauffman, Stuart, Jose Lobo, and William G. Macready. 2000. "Optimal Search on a Technology Landscape." *Journal of Economic Behaviour and Organization* 43 141-166.
- Nemet, Gregory F. 2012. "Inter-technology Knowledge Spillovers for Energy Technologies." *Energy Economics* 34 1259-1270.

- Nemet, Gregory F., and Evan Johnson. 2012. "Do Important Inventions Benefit from Knowledge Originating in Other Technological Domains?" *Research Policy* 41 190-200.
- Ridley, Matt. 2010. *The Rational Optimist*. New York, NY: HarperCollins Publishers.
- Schilling, Melissa A., and Elad Green. 2011. "Recombinant Search and Breakthrough Idea Generation: An Analysis of High Impact Papers in the Social Sciences." *Research Policy* 40 1321-1331.
- Schoenmakers, Wilfred and Geert Duysters. 2010. "The Technological Origins of Radical Inventions." *Research Policy* 39 1051-1059.
- Trajtenberg, Manuel. "A Penny for your Quotes: Patent Citations and the Value of Innovations." In *Patents, Citations and Innovation*, by Adam B. Jaffe and Manuel Trajtenberg, 25-50. Cambridge, MA: MIT Press, 2002.
- US Patent and Trademark Office. 2012. "Overview of the U.S. Patent Classification System (USPC)." *The United States Patent and Trademark Office*. 12. Accessed 10 15, 2014. <http://www.uspto.gov/patents/resources/classification/overview.pdf>.
- . 2014a. *Table of Issue Years and Patent Numbers, for Selected Document Types Issued Since 1836*. March 26. Accessed October 16, 2014. <http://www.uspto.gov/web/offices/ac/ido/oeip/taf/issuyear.htm>.
- . 2014b. "U.S. Manual of Classification File (CTAF)." *USPTO Data Sets*. Accessed August 2014. <http://patents.reedtech.com/classdata.php>.
- . 2014c. "U.S. Patent Grant Master Classification File (MCF)." *USPTO Data Sets*. Accessed August 2014. <http://patents.reedtech.com/classdata.php>.
- Weitzman, Martin L. 1998. "Recombinant Growth." *Quarterly Journal of Economics* 2 331-360.

Appendix – Robustness Checks for Section 7

This appendix contains the results from a series of robustness checks on the regressions discussed in section 7.3.

A1. Different Time Frames

As shown in Figure 1 in the main text, the fraction of patents with more than one mainline varies substantially over the period 1836-2012. Since the model is premised on ideas consisting of combinations of at least 2 elements, the model is a better fit for the data after 1936, when 69% of patents were assigned 2 or more mainlines, as opposed to 40% before 1936. This suggests the pre-1936 data may be subject to greater measurement error, which would lead to attenuation bias. Accordingly, I re-estimate my model for two time periods, 1836-1935 and 1936-2012. The results are presented in Table A1.

When I restrict attention to the period 1836-1935, the results are no longer statistically significant. However, for the second period, encompassing much of the 20th century, the results are significantly strengthened relative to the complete sample (compare to Table 7, column 4-6). This supports the argument that attenuation bias due to measurement error has likely reduced the size of the estimated coefficients.

As discussed in section 7.3, another potential source of bias in the estimates is the construction of the proxies. To check for the sensitivity of the results to the number of lags of the dependent variable, I double the number of lags in the last three columns of Table A1. This leads to a strengthening of the results, compared to those found in Table 7, column (4)-(6).

A2. Measuring Changes from the Prior Set

In Table A2, I measure the change in expected affinity for a class with reference to the pairs used in the previous period, rather than the current period. That is, when comparing Average Affinity Proxy in period t and period $t - 1$, I now use the set of pairs defined in $t - 1$. This has the effect of omitting all pairs that were used for the first time in period t , and restricts attention to pairs that have already been used at least once. As can be seen in the first three columns of Table A2, this has a significant impact on the coefficients attached to $\Delta \ln(\text{Average Affinity Proxy})_{cl,t-1}$. For all three proxies, some of the lagged changes now have a negative and statistically significant coefficient, in violation of Prediction 1.

Table A1: Robustness over time periods

Years Proxy	$\Delta \ln n_{cl,t}$								
	1844-1935			1944-2012			1852-2012		
	1	2	3	1	2	3	1	2	3
Age	-0.002*** (0.0002)	-0.002*** (0.0002)	-0.002*** (0.0002)	-0.0005*** (0.0001)	-0.0005*** (0.0001)	-0.0005*** (0.0001)	-0.001*** (0.0001)	-0.001*** (0.00005)	-0.001*** (0.00005)
$L^3 \Delta \ln(\text{Average Affinity Proxy})$	0.034 (0.067)	0.020 (0.057)	0.023 (0.047)	1.132*** (0.354)	0.899*** (0.273)	0.783*** (0.214)	0.282*** (0.076)	0.208*** (0.062)	0.214*** (0.050)
$L^4 \Delta \ln(\text{Average Affinity Proxy})$	0.061 (0.061)	0.055 (0.051)	0.055 (0.044)	0.168 (0.380)	0.169 (0.309)	0.434** (0.211)	0.129 (0.089)	0.111 (0.074)	0.165*** (0.058)
$L^5 \Delta \ln(\text{Average Affinity Proxy})$	-0.021 (0.060)	-0.013 (0.052)	0.030 (0.043)	0.145 (0.249)	0.114 (0.208)	0.361** (0.163)	0.245*** (0.079)	0.191*** (0.065)	0.206*** (0.050)
Controls									
Lags of $\Delta \ln n_{cl,t}$	8	8	8	8	8	8	16	16	16
Lags of $\Delta \ln N_t$	3	3	3	3	3	3	6	6	6
Class Fixed Effects	Y	Y	Y	Y	Y	Y	Y	Y	Y
Observations	23,766	23,766	23,766	27,173	27,173	27,173	50,155	50,155	50,155
Adjusted R ²	0.181	0.181	0.181	0.126	0.126	0.126	0.140	0.140	0.140

Note:

* $p < 0.1$; ** $p < 0.05$; *** $p < 0.01$
 White standard errors clustered by class in parentheses.

Table A2: Set of Pairs Defined by Earlier Period

Years Proxy	$\Delta \ln n_{cl,t}$								
	1844-2012			1844-1935			1944-2012		
	1	2	3	1	2	3	1	2	3
Age	-0.001*** (0.0001)	-0.001*** (0.00004)	-0.001*** (0.00005)	-0.003*** (0.0002)	-0.002*** (0.0002)	-0.002*** (0.0002)	-0.0005*** (0.0001)	-0.001*** (0.0001)	-0.001*** (0.0001)
$L^3 \Delta \ln(\text{Average Affinity Proxy})$	-0.264 (0.466)	-0.841*** (0.292)	-0.463*** (0.121)	-1.647*** (0.514)	-1.184*** (0.322)	-0.250* (0.128)	7.178*** (1.839)	2.947*** (0.996)	0.532 (0.565)
$L^4 \Delta \ln(\text{Average Affinity Proxy})$	0.655 (0.477)	-0.200 (0.320)	-0.274** (0.124)	-0.361 (0.522)	-0.370 (0.336)	-0.053 (0.133)	0.179 (2.492)	0.042 (1.243)	0.221 (0.441)
$L^5 \Delta \ln(\text{Average Affinity Proxy})$	-1.031** (0.465)	-1.092*** (0.302)	-0.068 (0.108)	-1.912*** (0.517)	-1.309*** (0.333)	0.126 (0.118)	-1.137 (1.850)	-0.446 (0.921)	0.592 (0.522)
Controls									
Lags of $\Delta \ln n_{cl,t}$	8	8	8	8	8	8	8	8	8
Lags of $\Delta \ln N_t$	3	3	3	3	3	3	3	3	3
Class Fixed Effects	Y	Y	Y	Y	Y	Y	Y	Y	Y
Observations	53,930	53,930	53,930	23,766	23,766	-0.002***	27,173	27,173	27,173
Adjusted R ²	0.145	0.146	0.146	0.183	0.184	(0.0002)	0.126	0.125	0.125

Note:

* $p < 0.1$; ** $p < 0.05$; *** $p < 0.01$
 White standard errors clustered by class in parentheses.

Further investigation reveals this result is not consistent over time. When I restrict our data to the first century, the negative coefficients remain statistically significant, and generally increase in magnitude (see the middle three columns in Table A2). However, when I restrict our attention to the 1936-2012 period (see the last three columns in Table A2), the sign on two of the proxies flips from negative to positive and statistically significant, and the coefficients associated with Proxy 3 are positive but insignificant. As noted above, I believe the results from the later period are more reliable, since the data better fits the model during this period.

Turning first to the later period, the change in coefficients is primarily due to the decrease in variation in $\Delta \ln(\text{Average Affinity Proxy})_{cl,t-l}$ when measured with the prior set of pairs. For example, $\Delta \ln(\text{Average Affinity Proxy 1})_{cl,t-l}$ has a mean of 0.005 when measured with the prior set of pairs, compared to 0.034 when measured with the most recent set of pairs. The means of $\Delta \ln(\text{Average Affinity Proxy 2})_{cl,t-l}$ and $\Delta \ln(\text{Average Affinity Proxy 3})_{cl,t-l}$ actually become negative, as well as closer to zero, when measured with the prior period's pair set. This is primarily because about half of pairs are only assigned to a patent once, and therefore much of the period-to-period change in Average Affinity Proxy comes from pairs being assigned their first and only patent. Such changes are only picked up by measuring $\Delta \ln(\text{Average Affinity Proxy})_{cl,t-l}$ with respect to the second period. Thereafter, these pairs either do not change (in Proxy 1), decline for one period (in Proxy 3), or decline by a small amount every period (in Proxy 3). While Proxy 3 appears incapable of picking up the positive impact of Average Affinity Proxy for the 1944-2012 period, Proxies 1 and 2 do, in support of the model.

As discussed in Section 8.3, a potential explanation for the negative coefficients in the earlier sample may be that the USPTO's updated classification has induced bias into the definition of mainlines.

A3. Alternative Sample Selection

In the first three columns of Table A3, I restore the $n_{cl,t} = 0$ observations by transforming the dependent variable to $\Delta \ln(1 + n_{cl,t})$. This leads to a weakening of the coefficients, but they remain positive and statistically significant. In the last three columns of Table A3, I include only classes with more than 1 mainline nested under the class, in case such classes are unusual or the mainlines they

draw on are poor proxies. This does not have a significant impact on estimated coefficients, since there appear to be few patents assigned to such classes.

Table A3: Alternative Samples

Dependent Variable:	$\Delta \ln(1 + n_{cl,t})$			$\Delta \ln n_{cl,t}$		
Omitted Classes	-			<2 Mainlines		
Proxy	1	2	3	1	2	3
Age	-0.001*** (0.00004)	-0.001*** (0.00004)	-0.001*** (0.00004)	-0.001*** (0.00005)	-0.001*** (0.00005)	-0.001*** (0.00005)
$L^3 \Delta \ln(\text{Average Affinity Proxy})$	0.080** (0.035)	0.065** (0.030)	0.058** (0.026)	0.210*** (0.062)	0.163*** (0.052)	0.155*** (0.043)
$L^4 \Delta \ln(\text{Average Affinity Proxy})$	0.163*** (0.033)	0.141*** (0.028)	0.136*** (0.025)	0.198*** (0.060)	0.168*** (0.050)	0.191*** (0.041)
$L^5 \Delta \ln(\text{Average Affinity Proxy})$	0.105*** (0.032)	0.100*** (0.027)	0.113*** (0.024)	0.078 (0.056)	0.067 (0.049)	0.113*** (0.041)
Controls						
Lags of $\Delta \ln n_{cl,t}$	8†	8†	8†	8	8	8
Lags of $\Delta \ln N_t$	3	3	3	3	3	3
Class Fixed Effects	Y	Y	Y	Y	Y	Y
Observations	61,678	61,678	61,678	52,358	52,358	52,358
Adjusted R ²	0.168	0.168	0.168	0.145	0.145	0.145

Note: † - Controls are lags of $\Delta \ln(1 + n_{cl,t})$. White standard errors clustered by class in parentheses.

* $p < 0.1$; ** $p < 0.05$; *** $p < 0.01$

CHAPTER 3

OPTIMAL RESEARCH STRATEGIES WHEN INNOVATION IS COMBINATORIAL

Matthew S. Clancy

Abstract

I develop a knowledge production function where new ideas are built from combinations of pre-existing elements. Parameters governing the connections between these elements stochastically determine whether a new combination yields a useful idea. Researchers are Bayesians who update their beliefs about the value of these parameters and thereby improve their selection of viable research projects. This approach provides a micro-foundations for knowledge spillovers, knowledge accumulation and fishing out effects. I use a combination of special cases and computer simulations to show that this model generates many stylized features of the research process. In particular, the optimal research strategy is a mix of harvesting the ideas that look best, given what researchers currently believe, and performing exploratory research in order to obtain better information about the unknown parameters. Initially, it may be optimal to perform relatively myopic, rather than exploratory research.

0. Introduction

This paper develops a new model of knowledge production, based on two premises.

First, all knowledge is composed of pre-existing parts. There is no *creatio ex nihilo*, wherein new ideas or technologies spring into existence fully formed from out of a void. Look deep enough and even the most creative ideas reveal themselves to be complex structures and arrangements of parts that were already there. These “parts” may be methods, techniques, concepts, mental models, designs, relationships, conventions, symbols, materials, facts, and so forth. What makes an idea new and creative is not *what* it is built from, but *what combination* of parts it is built from.

Second, knowledge is not *just* a combination of pre-existing parts. It is the connections between parts that matters, rather than the parts in and of themselves. Components interact, often in unexpected ways, to produce novel effects possessed by none of the components on their own. These effects may be desirable or undesirable, and learning how parts work together in concert or in opposition is crucial.

This paper will show how such a microfounded model of knowledge production, coupled with a simple model of a learning, innovating agent, is consistent with several stylized facts about the knowledge discovery process. For example;

- Ideas can build on each other, but can also be exhausted.
- There is a natural mechanism for knowledge to accumulate and spillover from one application to another.
- Mature knowledge sectors will be primarily characterized by applied, rather than basic, research.
- Incremental innovation will tend to peter out in the absence of radical innovations.

Some of this model's predictions are less intuitive, but I will argue they are consistent with actual patterns of knowledge discovery:

- The early days of a technological paradigm will often be dominated by applied, rather than basic, research.
- The absence of “moonshots” and other attempts at radical innovation is consistent with a healthy research sector.

The plan for this paper is as follows. After a review of some related literature (Section 1), I will first lay out the knowledge production function used in this paper, and embed this function in a simple model of a solitary agent conducting research (Section 2). I then discuss the optimal research strategy when there is no uncertainty about model parameters (Section 3). In general, however, parameter uncertainty is a key feature of this model, and the next section discusses the nature of researcher beliefs about model parameters (Section 4). I then discuss the optimal research strategy in the special case where the model takes the form of a multi-armed bandit problem (Section 5), before presenting the more complex general version of this model (Section 6). At this point, closed form solutions will be impossible, so I will numerically solve the problem for a large set of differing parameter values (Section 7). I will then characterize these numerically generated solutions (Section 8). Finally, I discuss the results and offer some thoughts on directions for further research (Section 9).

1. Background

The combinatorial nature of knowledge is clearest for physical objects like technology. Consider a laptop. Besides being an object with features and behaviors I value – for example, being able to run programs – it is also a configuration of tightly integrated components: a screen, trackpad, keyboard, hard drive, battery, etc. All of these components existed, in some form or another, before the invention of the laptop. In an important sense, the invention of the laptop consisted in finding a configuration of suitable components that could operate in concert. A similar exercise can be performed on other technologies. For example, a car is an integrated system of wheels, guidance systems, engines, structural supports, etc.

Non-physical creations can also be understood as combinations. Works of fiction draw on a common set of themes, styles, character archetypes, and other tropes; musical compositions rely on combinations of instruments, playing styles, and other conventions; and paintings deploy common techniques, symbols, and conventions. Indeed, even abstract ideas can be understood this way. In an essay on mathematical creation, Henri Poincaré, noted, “[Mathematical creation] consists precisely in not making useless combinations and in making those which are useful and which are only a small minority.”¹

Weitzman (1998) is the first to incorporate this feature of knowledge creation into the knowledge production function. In Weitzman’s model, innovation consists of pairing “idea-cultivars”² to see if they yield a fruitful innovation (a new idea-cultivar), where the probability an idea-pair will bear fruit is an increasing function of research effort. If successful, the new idea-cultivar is included in the set of possible idea-cultivars that can be paired in the next period. Weitzman’s main contribution is to show that combinatorial processes eventually grow at a rate faster than exponential growth processes, so that, absent some extreme assumptions about the cost of research, in the limit growth eventually becomes constrained by the share of income devoted to R&D rather than the supply of ideas. Simply put, combinatorial processes are so fecund that we will never run out of ideas, only the time needed to explore them all.

Weitzman’s model is echoed in Arthur (2009), who views all technologies as hierarchical combinations of sub-components. Arthur agrees that the laptop is a combination of screen, trackpad, keyboard, hard drive, battery, and so forth, but goes further, pointing out that, say, the

¹ Poincaré (1910), p. 325.

² So-called because the hybridization of ideas in the model is analogous to the hybridization of plant cultivars.

hard drive, is itself a combination of disks, a read/write head assembly, motors to spin and move these assemblies, and so forth. These sub-components themselves are combinations of still further subcomponents. Ridley (2010) also proposes a model akin to Weitzman's, arguing the best innovations emerge when "ideas have sex." This suggests ideas are combinations of genes which, when mixed, yield new offspring. This metaphor also anticipates the second premise, since it is the interaction of genes that matters in most instances.

Because clearly there is more to invention than simple combination. It does not suffice for engineers to bolt together elements at random, nor should musicians compose with a computer program that generates arbitrary lists of instruments, players and themes. Indeed, if we go back to the musings of Henri Poincaré, the full quote is:

In fact, what is mathematical creation? It does not consist in making new combinations with mathematical entities already known. Any one could do that, but the combinations so made would be infinite in number and most of them absolutely without interest. To create consists precisely in not making useless combinations and in making those which are useful and which are only a small minority. Invention is discernment, choice.

-Poincaré (1910), p. 324-325

Later Poincaré compares mathematical creation to the jostling of atoms, which are hoped will hook onto each other in stable configurations. He explains "The mobilized atoms are... not any atoms whatsoever; they are those from which we might reasonably expect the desired solution."³ In this paper I argue that we create by trying combinations similar in their composition to ideas with desirable features.

Several papers have explored this perspective, devising ways to measure the "distance" between ideas. Jovanovic and Rob (1990) represents a technology by an infinite vector, each element of which ranges between 0 and 1. The elements of this vector have an interpretation as methods, and the value of the element indexes how the method is used.⁴ Technologies are production functions and agents learn the mapping from technology vectors to productivity via Bayesian updating. Research consists in changing the values of the elements in a vector and observing the labor productivity associated with the new vector.

A related approach is developed by Kauffman, Lobo and Macready (2000) and Auerswald et al. (2000). These papers follow Jovanovic and Rob (1990) in thinking of technologies as a large

³ Poincaré (1910), p. 333-334.

⁴ For example, element k might be encoding "drill for oil at location k " and the value between 0 and 1 encodes some measure of the depth of drilling.

combination of distinct operations, although here the length of a technology vector is finite and each element can take on one of a finite number of states (rather than ranging over a continuous interval). The mapping between each technology vector and its productivity level is called a fitness landscape. When states are interdependent, the authors show this landscape is characterized by many local maxima. Innovation in such a model consists of exploring the fitness landscape by changing different operations. The roughly correlated nature of the landscape means small changes are likely to result in productivities that are similar to current levels, and large changes are essentially a draw from the unconditional distribution of productivity values. To reach the global maximum from any given position, it may be necessary to first traverse productivity “valleys.” Auerswald et al. (2000) uses the framework to show random deviations in a production process (akin to mutations in biology) can replicate many of the features of so-called “learning curves.”

The model that I develop in this paper combines the explicitly combinatorial framework of Weitzman (1998) with the vector based learning models of Jovanovic and Rob (1990), Kauffman, Lobo and Macready (2000) and Auerswald et al. (2000). This framework permits the derivation of many stylized facts about innovation to emerge from a common set of principles.

2. Model Basics

2.1. The Knowledge Production Function

I now describe formally how ideas are created in this model.

Definition 1: Primitive Elements. Let Q denote the set of primitive elements of knowledge q that can be combined with other elements to produce ideas, where $q \in Q$.

Definition 2: Pairs. Let p denote a two-element subset of Q , or “pair,” and P denote the set of two-element subsets of Q , where $p \in P$.

Definition 3: Ideas. An *idea* d is a set of pairs p , satisfying the condition that if $p_0 \in d$ and $p_1 \in d$, then $p \in d$ for any $p \subseteq p_0 \cup p_1$.

This model assumes a fixed set Q of technological building blocks that can be assembled into ideas, where any idea must contain at least two q from the set Q . For convenience, however, I define ideas in terms of the *pairs* of elements contained therein. So for example, an idea combining

elements q_1 , q_2 , and q_3 is represented as the set of subsets $((q_1, q_2), (q_1, q_3), (q_2, q_3))$. The condition attached to Definition 3 merely insures the pairs between all elements in the idea are included in the idea, so that we do not have ideas such as $((q_1, q_2), (q_1, q_3))$, which uses elements q_1 , q_2 , and q_3 but does not include the pair corresponding to (q_2, q_3) .

There are three important concepts in this model.

Definition 4: Compatibility. The *compatibility* of pair p in idea d is $c(p, d) \in \{0, 1\}$. When $c(p, d) = 1$ then the pair p is compatible in d . When $c(p, d) = 0$ then the pair p is incompatible in d .

Note that $c(p, d) = c(p, d')$ is not generally true. The compatibility of a pair may be equal to 1 in one idea and 0 in another.

Definition 5: Affinity. The probability a pair p is compatible defines its *affinity* $a(p) \in [0, 1]$.

The notions of compatibility and affinity are related as follows:

$$c(p, d) = \begin{cases} 1 & \text{with probability } a(p) \\ 0 & \text{with probability } 1 - a(p) \end{cases} \quad (1)$$

Essentially, this model assumes pairs of elements have an underlying tendency to be compatible or incompatible, and this tendency is described by the affinity of the pair.

Lastly, ideas are either *effective* or *ineffective*, where an idea is effective if and only if all the pairs of its constituent elements are compatible.

Definition 6: Efficacy. An idea d is *effective*, represented by $e(d) = 1$, iff $c(p, d) = 1 \forall p \in d$. In all other cases, represented by $e(d) = 0$, idea d is *ineffective*.

Restated, affinity determines the probability a pair is compatible, and when all pairs in an idea are compatible, the idea is effective. We may imagine ideas as sets of interacting elements that must be

mutually compatible for the idea to prove useful. If any two elements are incompatible, I assume the idea suffers a catastrophic failure that renders it unfit for use. Note the probability an idea is effective can be written as:

$$\Pr(e(d) = 1) = E[e(d)] = \prod_{p \in d} a(p) \quad (2)$$

This is the joint probability that every pair in the idea is compatible. Ideas are most likely to be effective when they are composed of elements that have a high affinity for each other, and least likely to be effective when composed of elements with a low affinity for each other.

Whereas I believe that this formulation for the structure of ideas strikes the right balance between simplicity and realism, a few caveats are in order, which I address here.

First, this model assumes every pair is equally important. In reality, technology is modular, with some elements tightly coupled and others only weakly interacting. A desktop computer, for example, consists of a monitor and the computer. The elements that make up the monitor are tightly interacting, but only loosely impacted by the elements in the computer. One way to capture this feature would be to assume, as in Weitzman (1998) that new combinations become elements available for combination in the future. However, as economists have long emphasized, unintended consequences are a prevalent feature of reality. Just because an inventor did not expect two elements to interact with each other does not mean they will not do so in unanticipated ways. Suppose the probability that two elements are incompatible proceeds in two steps. First, there is some probability that the two elements interact with each other. Then, if they do, there is a second probability that they interact poorly, causing the entire assemblage to break down. This is a perfectly valid way to understand what is being captured by the single parameter affinity.

Second, the model assumes that only pairwise interactions between the building blocks of ideas matter. There are obvious counter-examples. Suppose elements q_1 and q_2 are chemical compounds that do not react unless in the presence of a catalyst q_3 . Or imagine that a stabilizer regulates the interaction between two components that would otherwise interact in a calamitous manner. Higher order interactions are equally plausible. One impact of allowing for interactions above the pair level would be to make learning more difficult, since observing the behavior of a single pair of elements is less informative if the presence or absence of a third element is crucial. Moreover, if we determined, for example, that only interactions between sets of three elements matter, many of the results could

be derived with appropriate redefinitions (for example, affinity would now apply to sets of three elements).

2.2. The Conduct of Research

This production function is used by a researcher who is trying to discover effective ideas. The researcher is a risk-neutral infinitely-lived profit maximizer with a discount factor $\delta \in (0,1)$. She knows every element in the set Q , and in each period may choose to conduct a research project on some idea d built from the elements in Q .

Definition 7: Possible Ideas. The set D_P is the set of all possible ideas that can be made from elements in Q . It contains all subsets of P that satisfy the condition in Definition 3.

Definition 8: Eligible Ideas. A set of eligible ideas $\tilde{D}$ is a subset of D_P . It is only sensible to conduct research projects on eligible ideas, and when a research project is attempted, the idea is removed from $\tilde{D}$ at the end of the period.

The set $\tilde{D}$ is primarily intended to indicate the set of *untried* ideas, and so it shrinks as research proceeds. I add to this set an additional element, the null set $d_0 \equiv \{\emptyset\}$, which represents the option not to conduct research in a period.

Definition 9: Available Actions. The researcher's set of available actions is $D \equiv d_0 \cup \tilde{D}$.

Note that because $d_0 \notin \tilde{D}$, if the researcher chooses not to conduct research, then this option is not removed from its action set in the next period.

In principle, the researcher “knows” every idea that can be built from elements in Q , in the same sense that I “know” every economics article that can be written with words and symbols in my repertoire. However, just as I do not know whether any of these articles are good until I think more about them, or actually write them out, the researcher does not learn if an idea is effective until she

decides to conduct research on it.⁵ Indeed, research is costly, requiring investments of time and other resources. I assume that research on any idea has cost $k(d)$, known to the researcher, and that the option d_0 , to do nothing, has $k(d_0) = 0$.

The return from conducting research is a reward $\pi(d)$, also known to the researcher, which is received if the idea is eligible and discovered to be effective. This reward could indicate a prize for innovation from the government, or the sale of patent rights over the idea to a firm, or some other incentive for innovation. Because each chosen idea is removed from the set of eligible ideas D at the end of a period, the researcher cannot claim a reward for the same idea multiple times. I assume the reward value of the outside option d_0 is always zero.

Hence, a researcher who chooses to conduct research on idea d expects to receive a net value of:

$$\pi(d)E[e(d)] - k(d) \quad (3)$$

The idea is successful with probability equal to the expected efficacy $E[e(d)]$, in which case the researcher obtains a reward $\pi(d)$. Whether the idea succeeds or not, the researcher pays up front research costs $k(d)$. This formulation of the innovator's problem is not unusual, except for the term $E[e(d)]$, which is determined by the knowledge production function described earlier.

Recall that $e(d) = 1$ (an idea is effective) if and only if all of its pairs are compatible. This implies the probability distribution of $e(d)$ is:

$$e(d) = \begin{cases} 1 & \text{with probability } \prod_{p \in d} a(p) \\ 0 & \text{with probability } 1 - \prod_{p \in d} a(p) \end{cases} \quad (4)$$

Therefore:

⁵ Jorge Luis Borges tells a parable of an infinite library containing books with every combination of letter and punctuation mark. In this library, there is a book resolving the basic mysteries of humanity, since every possible book exists, but finding the book and verifying it is true amongst all the gibberish and babel is a daunting task for the library inhabitants. See Borges (1962).

$$E[e(d)] = \prod_{p \in d} a(p) \quad (5)$$

Hence, if the researcher knows the affinity of each pair, she can compute the expected efficacy of every idea. In general, I will assume the researcher does *not* know the true affinity of each pair and must infer its likely value from the outcomes of research projects.

Before proceeding to this more complex and realistic case, I discuss the special case where the agent knows the true affinity of each pair with certainty.

3. Special Case 1: Affinity is Known

Suppose there is a researcher with perfect knowledge of $a(p)$. In any period, the researcher's problem is to determine which idea, if any, to attempt. Define the expected present discounted value of the optimal strategy in period t as:

$$V_t = \sum_{\tau=0}^{\infty} \delta^{\tau} (\pi(d_{t+\tau}) E[e(d_{t+\tau})] - k(d_{t+\tau})) \quad (6)$$

where d_t denotes the optimal decision in period t . Note that, because no relevant information is revealed by the outcome of research, the optimal choice in period $t + \tau$ depends only on information available in period t . Recall also that researchers can always choose d_0 , to obtain a payoff of zero with certainty. Because the set of eligible ideas is finite, the researcher will always resort to choosing d_0 in the end.

The optimal strategy in the absence of learning is simple:

Remark 1: Optimal strategy with certainty. In each period, the optimal strategy is to choose the eligible idea with the highest $E[e(d)]\pi(d) - k(d)$ so long as $E[e(d)]\pi(d) - k(d) \geq 0$. If no idea satisfies this, then choose d_0 .

Because there is no learning in this special case, the researcher's problem collapses to choosing the order in which to consume a set of lotteries. Because the researcher is risk-neutral and discounts the future, she orders these lotteries in descending expected (net) value. Furthermore, the researcher

never attempts a lottery where cost exceeds expected value. Because there is no learning, and because the researcher uses up the best lotteries first, the following remark holds.

Remark 2: Value decreases over time. The anticipated value of research declines over time, i.e.,

$$V_t \geq V_{t+\tau} \quad \forall t, \tau > 0 \quad (7)$$

In the absence of learning, the best research ideas are used up (“fished out”) first, so that the value of conducting research falls over time. This stands in contrast to the prevalent view that knowledge is cumulative, and that today’s researchers can accomplish more than their forebears by building on their accomplishments (“standing on the shoulders of giants”). This result is not general, of course, but a consequence of the absence of learning.

In general, the researcher does *not* know the affinity of a pair with certainty. However, the problem faced by a researcher becomes close to this limiting special case when they already possess very good information about most pairs, so that firms are *nearly* certain about the true affinity of each pair. Thus, mature technology sectors should be more characterized by myopic research strategies than young ones. Alternatively, it may be difficult for firms to capitalize on information they learn, if the outcomes of research is difficult to conceal from other firms. If there are many firms operating in a sector, and if knowledge diffuses rapidly, then competitors can expropriate the value of information. If it is difficult to conceal information (for example, if discoveries are products that can be easily reverse engineered), then firms may behave more or less myopically, even if the sector is young. Finally, firms may also behave myopically if learning is very costly or difficult. In any case, a discussion of firm learning is necessary. I turn to this topic next.

4. Research and Learning

In general, I assume the affinity $a(p)$ between a pair of elements is *ex ante* unknown to researchers. Instead, researchers are Bayesians with prior beliefs over the possible distribution of $a(p)$. Though the researcher does not observe $a(p)$ directly, she can make educated guesses based on the tendency of p to be compatible or incompatible. Using her beliefs about the affinities between all pairs in an idea, she can compute the probability an idea will be effective. In more formal terms, a crucial part of the discovery process is the inference of likely affinity values from the compatibility or incompatibility of component-pair interactions.

I impose the following assumption on the researcher's beliefs:

Assumption 1: Independence of Affinity. The researcher believes $a(p)$ is independently distributed for all p .

As long as this assumption stands, the updating of beliefs about any $a(p)$ depends only on observations on the pair p alone. If $a(p)$ were not independently distributed, it would be necessary to also take into account the observations on correlated pairs, greatly complicating the problem.

Each observation of compatibility is the outcome of a Bernoulli trial governed by the pair's true affinity, with the two possible states being compatibility (probability $a(p)$) or incompatibility (probability $1 - a(p)$). Given s instances of compatibility ("success") and f instance of incompatibility ("failure"), the researcher updates her beliefs according to Bayes law under the Bernoulli distribution:

$$\Pr(a(p) = \tilde{a} \mid s, f) = \binom{s+f}{s} \tilde{a}^s (1 - \tilde{a})^f \frac{\Pr(a(p) = \tilde{a})}{\int_0^1 \binom{s+f}{s} a^s (1 - a)^f \Pr(a(p) = a) da} \quad (8)$$

where $\binom{s+f}{s} \tilde{a}^s (1 - \tilde{a})^f = \Pr(s, f \mid a(p) = \tilde{a})$.

In the remainder of the paper, I impose the following assumption:

Assumption 2: Beta Distributions. The researcher believes $a(p)$ follows a beta distribution.

The beta distribution has the useful property of being the conjugate family of a Bernoulli distribution. The conjugate family for a distribution defines a class of distributions such that, if the prior distribution belongs to the conjugate family, then the posterior distribution will as well. In this case, if the prior distribution for an affinity belongs to the beta distribution, then after updating the researcher's beliefs by observing a set of Bernoulli trials, the posterior distribution will also be a (different) beta distribution. This structure maintains a constant form for the beliefs of the

researcher as her information varies. Moreover, the form of a beta-distribution is sufficiently flexible to enable the exploration of many kinds of assumptions about the prior beliefs of the researcher.

The probability density function of a beta distribution takes the following form:

$$f(a) = \frac{\Gamma(\alpha + \beta)}{\Gamma(\alpha)\Gamma(\beta)} a^{\alpha-1} (1-a)^{\beta-1} \quad (9)$$

where $\Gamma(x)$ is the gamma function. The distribution's support is over the $[0,1]$ interval with its shape governed by the parameters $\alpha > 0$ and $\beta > 0$. Changing α and β can yield a centered bell shape, highly skewed distributions, and U-shaped distributions. It can be shown that, given a beta distribution:⁶

$$\int_0^1 \binom{n}{s} a^s (1-a)^{n-s} f(a) da = \binom{n}{s} \frac{\Gamma(\alpha + \beta)}{\Gamma(\alpha)\Gamma(\beta)} \frac{\Gamma(s + \alpha)\Gamma(n - s + \beta)}{\Gamma(n + \alpha + \beta)} \quad (10)$$

Combining equations (8), (9), and (10), the updated beliefs of the researcher given s instances of compatibility from n total observations is:

$$\Pr(a(x) = \tilde{a} | s, n) = \frac{\Gamma(n + \alpha + \beta)}{\Gamma(s + \alpha)\Gamma(n - s + \beta)} \tilde{a}^{s+\alpha-1} (1 - \tilde{a})^{n-s+\beta-1} \quad (11)$$

Note that this is equivalent to a beta distribution with $\alpha' = s + \alpha$ and $\beta' = n - s + \beta$. Hence, defining $\alpha(p)$ and $\beta(p)$ to be the initial parameters governing the prior beliefs about of $a(p)$, after observing $s(p)$ instances of compatibility and $n(p) - s(p)$ instances of incompatibility, the researcher believes $a(p)$ to be governed by a beta distribution with parameters $\alpha(p) + s(p)$ and $\beta(p) + n(p) - s(p)$. The expected value of such a distribution⁷ is given by:

$$E[a(p) | n(p), s(p)] = \frac{\alpha(p) + s(p)}{\alpha(p) + \beta(p) + n(p)} \quad (12)$$

Note that as the number of observations grows large, the expectation converges to $s(p) / n(p)$, which is simply the proportion of observations where compatibility is observed. The sum

⁶ See Casella and Berger (2002) p. 325.

⁷ See Casella and Berger (2002) p. 325.

$\alpha(p) + \beta(p)$ determines the relative weight put on new observations and the initial beliefs, and is a measure of initial certainty.

When a researcher knows the affinity of all pairs with certainty:

$$E[e(d)] = \prod_{p \in d} a(p) \quad (13)$$

and given Assumption 1, this can be expressed as:

$$E[e(d)] = \prod_{p \in d} E[a(p)] \quad (14)$$

where $E[a(p)]$ is given by equation (12). This expression captures the core notion of this model: ideas are more likely to be effective (“useful”) when composed of elements that have worked well together frequently in the past.

Under some special assumptions, this learning framework, coupled with our model of innovation, can be modeled as a multi-armed bandit problem, which I discuss in the next section.

5. Special Case 2: Research as a Multi-armed Bandit Problem

To isolate the effects of learning in the researcher’s problem, I reformulate the problem as a multi-armed Bernoulli bandit problem. While this imposes strong restrictions on the model, the advantage is that there exists a large literature on such problems.⁸ In such problems an agent must choose between n options, each of which offers a reward with a fixed probability unknown to the agent. Over time, as the researcher observes the frequency with which she receives a reward from any given option, she obtains a progressively better estimate of the underlying probability of receiving a reward from that option. The agent’s problem is to balance a myopic strategy that selects the option currently believed to be most favorable, and a far-sighted strategy which seeks to gather information on other options so that the true best option can be found. Essentially, she must make a trade-off between exploitation and exploration.

Such models have a well-known solution technique, called a Gittins index (discussed below). My model takes the form of a multi-armed Bernoulli bandit problem under the following conditions:

1. There are n pairs p_i , where $i = 1, \dots, n$. These pairs have no elements in common.

⁸ See Gittins, Glazebrook, and Weber (2011).

2. The researcher correctly believes $a(p_i)$ follows a beta distribution with parameters α_i and β_i .
3. Each pair may be combined with any number of element from the set Q , which is infinitely large. These ideas:
 - a. Have prize value $\pi(d) = 1$.
 - b. Have cost $k(d) = k$.
 - c. Have expected efficacy $E[e(d)] = E[a(p_i)]$ (which implies the researcher knows the other affinities are equal to 1 with certainty).
4. No other ideas are eligible.

In this setting, the researcher's problem collapses to choosing which pair p_i to combine with elements in Q (the identity of the elements does not matter, since they all have the same prize value, cost, and expected affinities). The researcher must balance the expected net reward:

$$E[a(p_i)] - k \quad (15)$$

against the value of learning more accurately the true value of each $a(p_i)$. Because the researcher believes all affinities except for $a(p_i)$ are equal to one, any time the attempted idea is effective, the researcher learns the pair p_i is compatible, and vice-versa.

Specifically, the researcher's problem is characterized by a Bellman equation of the form:

$$V(B) = \max \left\{ \max_{p_i} \left[\frac{\alpha_i}{\alpha_i + \beta_i} \{1 + \delta V(B_{i+})\} + \frac{\beta_i}{\alpha_i + \beta_i} \delta V(B_{i-}) - k \right], 0 \right\} \quad (16)$$

where

$$\begin{aligned} B &= ((\alpha_1, \beta_1), \dots, (\alpha_i, \beta_i), \dots, (\alpha_n, \beta_n)) \\ B_{i+} &= ((\alpha_1, \beta_1), \dots, (\alpha_i + 1, \beta_i), \dots, (\alpha_n, \beta_n)) \\ B_{i-} &= ((\alpha_1, \beta_1), \dots, (\alpha_i, \beta_i + 1), \dots, (\alpha_n, \beta_n)) \end{aligned} \quad (17)$$

Taking this equation from left to right, the researcher's problem is first to choose whether or not to conduct research. If she does not, she obtains 0 with certainty. Moreover, if the researcher ever

chooses to quit research in one period, she will do so in all subsequent periods, since her information and action set will be unchanged in the following period.

If the researcher does choose to conduct research, she must decide the best pair p_i to select. With probability $\alpha_i / (\alpha_i + \beta_i)$, the idea is effective and the researcher obtains a reward equal to 1. Moreover, in the next period, she will update her beliefs in accordance with equation (12), so that her beliefs are described by the vector B_{i+} . Therefore, in the next period, she obtains $V(B_{i+})$, discounted by δ . Conversely, with probability $\beta_i / (\alpha_i + \beta_i)$ the idea is ineffective and she obtains no reward this period. Moreover, if an idea is ineffective, the researcher learns pair p_i is incompatible and the researcher will update her beliefs to vector B_{i-} . In the next period, she will obtain $V(B_{i-})$, discounted by δ . Finally, in either case, she pays k to conduct research.

This formulation is equivalent to a standard multi-armed bandit problem with expected affinity of each pair corresponding to the fixed Bernoulli probability of receiving a reward. Using the Gittins Index approach, we can obtain some clear insights, summarized in the following remarks.

Remark 3: Optimal strategy with learning. The optimal strategy in every period is to choose the option with the highest Gittins Index $\lambda_i(\alpha_i, \beta_i)$ where

$$\lambda_i(\alpha_i, \beta_i) \equiv \sup \{ \lambda_i : v_i(\alpha_i, \beta_i, \lambda_i) = 0 \} \quad (18)$$

and $v_i(\alpha_i, \beta_i, \lambda_i)$ is equal to:

$$\max \left\{ \frac{\alpha_i}{\alpha_i + \beta_i} (1 + \delta v_i(\alpha_i + 1, \beta_i, \lambda_i)) + \frac{\beta_i}{\alpha_i + \beta_i} \delta v_i(\alpha_i, \beta_i + 1, \lambda_i) - k - \lambda_i, 0 \right\} \quad (19)$$

A proof is presented in the appendix. See Gittins, Glazebrook, and Weber (2011) for more discussion.

Note that $v(\alpha_i, \beta_i, \lambda_i)$ only depends on the beta parameters of one pair, p_i . This decomposes the n -dimensional choice problem into n one-dimensional problems. The Gittins index $\lambda_i(\alpha_i, \beta_i)$ can be thought of as a riskless payment the researcher can receive in lieu of the reward from choosing p_i , chosen so the researcher is exactly indifferent between the two options. It accounts for

the expected value in this period, equal to $\alpha_i / (\alpha_i + \beta_i) - k$, plus the prospects of achieving a better or worse outcome in subsequent periods, as the researcher's beliefs are updated. Choosing the highest Gittins index in every period therefore accounts for the rewards in the current period, plus the potential gains from better information.

Note, however, that the researcher's expected payoff is not equal to the Gittins index. This is because a Gittins index is computed with reference only to one pair. The index tells the researcher what to do, but it does not say how much she should expect to make. However, the value of research still follows some principles:

Remark 4: Value rises with affinity (value of learning). The expected value of research is nondecreasing in α_i .

In this model, value only comes from obtaining rewards, which occur with probability $\alpha_i / (\alpha_i + \beta_i)$. If the researcher fixes an optimal strategy, and then one α_i is increased, the researcher cannot be worse off. Neither will she be worse off if we let her re-select the optimal strategy.

Multi-armed bandit problems also have the following feature:

Remark 5: Stick with the winner. The optimal strategy follows a “stick-with-the-winner” formulation, i.e.,

$$\lambda_i(\alpha_i + 1, \beta_i) > \lambda_i(\alpha_i, \beta_i) > \lambda_i(\alpha_i, \beta_i + 1) \quad (20)$$

A proof is presented in Bellman (1956).

Once an idea has been found successful, in this model, the probability it will be successful in the next period as well is increased, which makes the choice still more favorable in the next period and the opportunity cost of trying something else higher. Therefore, the researcher always sticks with a winning pair, at least until it stops working (although this is not sufficient for her to switch his strategy either).

Besides preferring winners, the researcher also prefers the pair about which less is known:

Remark 6: Favor uncertainty. When the expected reward is the same, an optimal strategy chooses the option where more is learned, i.e.,

$$\lambda_i(\alpha_i, \beta_i) > \lambda_i(m\alpha_i, m\beta_i) \quad (21)$$

if $m > 1$.

A proof is presented in Gittins and Wang (1992).

Note that the expected value of a beta distributed variable with parameters (α, β) and $(m\alpha, m\beta)$ is the same, since:

$$\frac{m\alpha}{m\alpha + m\beta} = \frac{\alpha}{\alpha + \beta} \quad (22)$$

However, a pair with $(\alpha + 1, \beta)$ has a higher expected value than one with $(m\alpha + 1, m\beta)$. In the next period, the potential benefits are higher for the more uncertain idea. At the same time, the potential downsides also looms larger for the more uncertain choice. However, since the researcher always has the option to quit research, downside risks are capped at 0. This leads agents to prefer ideas from which they can learn more, that is ideas where they have less certainty about the affinity of their pairs.

Finally, the above results imply the following.

Remark 7: Value is rising and concave in success: The expected value of research is increasing in the number of times any given pair is found to be effective, denoted $s(p)$, and bounded from above by $(1 - k) / (1 - \delta)$.

That $V(B)$ is increasing in $s(p)$ is simply a reformulation of Remark 4, since the researcher's α_i parameter is updated to $\alpha_i + s(p)$ after observing $s(p)$ instances of compatibility. Moreover, since researchers stick with winners, as $s(p)$ continues to increase, the expected payoff begins to resemble the payoff from simply playing the same pair in each period. This is a concave function of $s(p)$, bounded from above by $(1 - k) / (1 - \delta)$.

To summarize, in the multi-armed bandit formulation of the researcher's problem, the optimal strategy is a mix of myopic and far-seeing strategies, since the Gittins index takes into account both the immediate payoff and the distant future payoffs. Researchers, somewhat paradoxically, prefer both proven research paths (they stick with winners) and unproven *and* untested research projects (they favor uncertainty). The value of research increases when research is successful, so that research is cumulative and has a standing-on-the-shoulders-of-giants effect. However, eventually, the payoff from success stops increasing as it approaches a ceiling, given by the present discounted value of a successful research project in every period.

However, to derive these results, I had to rely on some extreme assumptions that simplified the combinatorial dynamics of this model. This simplified special case is nonetheless a close approximation of the researcher's problem in some settings. It bears similarities to the literature on general purpose technologies. General purpose technologies are those like steam power, electricity, or computers, for which the new technology has many applications (Helpman 1998). These applications are well understood, and can be relatively easily developed, if the general purpose technology is developed. In the terms of this model, the (potential) new general purpose technology is like the pairs whose affinity is unknown, and the applications are a pool of technological building blocks where the affinity between them is very high, but where the technologies built from them alone have either been exhausted or are of low value.

Imagine, for example, several different ways to generate electrical power. If a successful platform for generating electrical power can be discovered, the researcher is confident it can be easily combined with a functionally limitless set of existing equipment. The different ways of producing electricity can be modeled as different pairs of elements (say, magnets and wires), each of which has an unknown affinity. The different technologies awaiting a new power source can be modeled as sets drawn from a set Q . Such a problem is very similar to the multi-armed bandit approximation discussed here, and the solution concept would be similar. In particular, researchers could proceed by independently weighing the prospects of each alternative method of generating electricity, as in the Gittins index approach. If the first approach turned out to work all the time, the researcher would never bother trying others. The value of her research agenda would asymptote at its maximum, with the researcher successfully extending her electrical system to a new application in every period. If the approach fails, however, in choosing the next method, she will put more weight on methods that are comparatively less well understood.

In the next section, I will incorporate uncertainty about the true affinity into a more general setting.

6. The General Case

6.1. Learning in the General Setting

To implement the Bayesian belief updating presented in section 4, I summarize the researcher's beliefs by the vector B where:

$$B = \left((\alpha_1, \beta_1), \dots, (\alpha_i, \beta_i), \dots, (\alpha_{m(N)}, \beta_{m(N)}) \right) \quad (23)$$

and $m(n) \equiv n(n-1)/2$. These beliefs are now updated via a stochastic vector $\omega(d)$ of information revealed by a research project over idea d . This vector has the same number of elements as B and is defined so that after conducting a research project on d , the updated beliefs vector B' is given by:

$$B' = B + \omega(d) \quad (24)$$

For example, if a research project on some d' reveals pair p_1 is compatible and pair $p_{m(n)}$ is incompatible, and does not reveal any other information, then $\omega(d')$ takes the form:

$$\omega(d') = ((1,0), (0,0), \dots, (0,0), (0,1)) \quad (25)$$

In Section 5, I restricted attention to research projects where $E[e(d)] = E[a(p_i)]$. In this setting, it was obvious that a research project would reveal the compatibility $c(p_i, d)$: since the only uncertainty pertained to this pair, if the idea was effective, it must have been that $c(p_i, d) = 1$, and if the idea was ineffective, it must have been that $c(p_i, d) = 0$.

In the general setting, where *every* pair in an idea may be compatible or incompatible, what is learned from research is more complicated. Whenever an idea is effective, it must be that $c(p, d) = 1$ for all $p \in d$ (since this is the definition of efficacy). Accordingly, whenever an idea is effective, the researcher must know that every $p \in d$ was compatible and update her beliefs accordingly.

However, whenever an idea is ineffective, any configurations of compatibilities where at least one $c(p, d) = 0$ is capable of generating the same result. Which pair compatibilites are revealed in

this case is not a trivial matter, and can easily make the model intractable or embed unrealistic features. The following procedure has two main virtues, discussed more below. First, I believe it has realistic features about what can be learned from successes and failures. Second, it keeps beliefs independent, which maintains the model's tractability.

The information about the compatibilities $c(p, d)$ of each pair $p \in d$, is generated by the following stochastic process.

Assumption 3: Learning From Research (Frustrations of Failure). The information revealed about the compatibility of pairs is determined according to the following procedure:

1. A pair $p \in d$ is randomly drawn with equal probability from among the pairs whose compatibility has not already been selected.
2. The compatibility of this pair is added to the research project's revealed information.
3. If $c(p, d) = 1$ and unselected pairs remain, return to step one and repeat the above procedure. If $c(p, d) = 0$ or if no unselected pairs remain, do not add any more compatibilities to the research project's revealed information.

When this procedure is completed, the researcher observes a packet of information $\omega(d)$. If an idea is effective, this revelation procedure will reveal that all of its pairs are compatible. If the idea is ineffective, it will reveal some, but possibly not all, of the compatibilities of the pairs that make up d . Specifically, the above revelation procedure will never reveal more than one pair is incompatible.⁹

This revelation mechanism is meant to capture the frustrations of failure. Researchers often have some indication of where things began to go wrong – for example, a proof step that does not go through, or an engine part that overheats – but a full understanding of why the idea failed is often elusive. This partial revelation of information about what part of the innovation failed is captured by the fragmentary knowledge of which pairs are compatible and incompatible when an idea is ineffective. If information is not fragmentary when an idea fails, then researchers could in principle

⁹ This assumption could be made more general with the introduction of a parameter $\eta \in [0, 1]$, so that when a pair with $c(p, d) = 0$ is revealed, the revelation procedure stops with probability η . The model presented here is then the special case with $\eta = 1$.

learn about the affinity between every idea simply by performing research on an idea that draws on every element in the set Q .

Note also the researcher is unlikely to learn much from a research project with many incompatible pairs, because the probability of encountering an incompatibility that stops the revelation process early is high. This is meant to capture the notion that we do not, on average, learn as much from a project that is wrong on many levels. When an idea has only one or two incompatible pairs, the researcher may learn a lot or a little by trying the research project. If she is lucky, it is the kind of idea where, although the idea is ineffective, she can get a long way before hitting a roadblock. Such a research project might be represented by one where the first incompatible pair is only reached after a long series of compatible ones. If she is unlucky, the idea is the kind in which it is very difficult to make any headway until a certain problem is cracked. This would be represented by an idea where the revelation of compatibilities is quickly stopped by an incompatible pair.

Lastly, not that $\omega(d_0) = ((0,0), \dots, (0,0))$ by assumption: agents learn nothing when they choose not to do a research project.

6.2. The Researcher's Problem in the General Setting

In this paper I limit attention to the case where the researcher has no competition, thereby evading strategic considerations. Hence the researcher's problem is to find an optimal policy function $d^*(D, B)$ mapping from available actions D and beliefs B to an optimal choice.

The policy function $d^*(D, B)$ maximizes:

$$V(D, B) = E \left[\sum_{t=0}^{\infty} \delta^t \left(e(d^*(D_t, B_t)) \pi(d^*(D_t, B_t)) - k(d^*(D_t, B_t)) \right) \right] \quad (26)$$

where

$$\begin{aligned} D_{t+1} &= d_0 \cup (D_t \setminus d^*(D_t, B_t)) \\ B_{t+1} &= B_t + \omega(d^*(D_t, B_t)) \end{aligned} \quad (27)$$

The payoff function in (26) is the discounted sum of expected per-period returns from choosing idea $d^*(D_t, B_t)$ in period t . The researcher obtains $\pi(d^*(D_t, B_t))$ if $d^*(D_t, B_t)$ is effective, which

occurs with probability $E[e(d^*(D_t, B_t))]$, but pays $k(d^*(D_t, B_t))$ either way. As noted earlier, if the researcher chooses d_0 , then $\pi(d_0) = E[e(d_0)] = k(d_0) = 0$.

It is instructive to write equation (26) as a Bellman equation:

$$V(D, B) = \max_{d \in D} \{E[e(d)]\pi(d) - k(d) + \delta E[V(D', B')]\} \quad (28)$$

where

$$\begin{aligned} D' &= d_0 \cup (D \setminus d) \\ B' &= B + \omega(d) \end{aligned} \quad (29)$$

This formulation makes clear that the choice of idea has a payoff in the current period, but also an impact on the future, through the $E[V(D', B')]$ term. Agents may prefer to take a loss in the current period, in order to learn and increase their payoffs in the future. This formulation also makes clear that if it is ever optimal for the researcher to choose d_0 at some stage, then it is optimal for her to do so in every subsequent period because such a choice ensures $D' = D$ and $B' = B$.

7. Numerical Simulation

7.1. Solution Methodology

There is no general closed-form solution for this problem, but it can (in principle) be solved by backwards induction. Because the set of possible ideas is finite, and (as noted above) the researcher never pauses then restarts research, it is certain that from period $|D|$ on the researcher will choose d_0 in every stage. Thus, $V(D, B) = 0$ in period $|D|$, no matter what the researcher does in period $|D| - 1$. Since the researcher knows the next period payoff with certainty, she can find the best choice in period $|D| - 1$, and work backwards toward period 0.

7.2. Simulation Assumptions

Although a closed form solution is not feasible, I would like to highlight some common contours of an optimal solution. Unfortunately, this model is also beset by the “curse of

dimensionality,” so that even a general numerical approximation is difficult to obtain.¹⁰ The curse of dimensionality is a name applied to problems where the computational resources needed to solve them grow very quickly.¹¹ For example, consider the number of parameters needed to characterize the state-space for a problem with 3, 4, and 5 elements. Given 3 elements, there are 3 pairs and 3 affinities, the beliefs about which are described by 2 parameters each, so that B has six dimensions. Three elements also implies there are 4 possible ideas. Since each of these can be available or unavailable, there are 2^4 different sets D associated with every B . If I increase the number of elements to 4, then B has 12 dimensions and there are 2^{11} possible sets of D for each vector of beliefs. If I increase the number of elements to 5, then B has 20 dimensions and there are 2^{26} different sets D associated with each one. Obtaining a good approximation of a state space with 20 dimensions and 2^{26} discrete action states is very challenging.

Given the foregoing, my approach is to instead solve a manageable (small) version of the problem 100 times, and then to use a regression analysis to see if characteristics of Remarks 1-7 hold for these optimal solutions. Essentially, I will project a linear approximation onto a highly complex and non-linear solution, to check for the validity of Remarks 1-7 outside of the special cases for which they were derived. It will be important that the choice of problem to solve is sufficiently rich that it captures the complexity of the model, but remains solvable. The basic structure of the problem I will solve takes the following form:

1. There are four elements $q \in Q$ that can be combined into ideas.
2. There are ten eligible ideas. Six ideas are composed of pairs of elements, and four ideas are composed of three elements.¹²
3. The cost of ideas is normalized to 1.
4. The researcher’s discount factor is 0.95.
5. The researcher’s beliefs about the affinity of each pair follow a beta distribution.

Each pair has unique beta parameters. I discuss this more below.

¹⁰ Of course, if a numerical simulation is difficult to achieve for a computer, then it is likely to be even more difficult for a human brain. The following section may also be interpreted as one set of heuristics that computationally-constrained researchers might use to approximate the optimal research strategy.

¹¹ Powell (2011) provides a good overview of the problem and possible approximation methods.

¹² Including the 11th idea, composed of all four elements, dramatically increases the computational time to solve, without adding much insight, so I omit it. Alternatively, we might assume this idea has prohibitively high costs, but is theoretically eligible.

6. Each idea has a unique prize value $\pi(d)$, known to the researcher. I discuss this more below.

These six conditions capture several key features. First, I model the beliefs of the researcher (as in Special Case 2), while allowing ideas to depend on multiple pairs and to have different reward values (as in Special Case 1). Second, each pair can be used up to three times (once in a two-element idea, and twice in a three-element idea), so that updating of beliefs can happen more than once. Third, information revealed from one idea spills over to other ideas. Fourth, some ideas are subsets of others.

Most importantly, this problem can be solved in a reasonable amount of time. I have written a computer program in python to solve this problem. To begin, I define the possible sets D in which the researcher may find herself. Since there are 10 ideas, and each idea may be either eligible or ineligible, there are $2^{10} = 1,024$ distinct sets of eligible ideas. For each of these sets, I define the potential belief vectors of interest. Since I know the initial beta parameters of every $a(p)$, the program can exhaustively list the vectors B that might be attained in any given set.

Once I have a set of (D, B) states, I work backwards. The program begins by evaluating the null set $D = \{\emptyset\}$, where the only available option is to quit research and earn zero with certainty (for every belief vector B). Next, using this result, it evaluates the best action for each set with just one eligible idea remaining. When there is just one eligible idea, the problem simplifies to:

$$V(D, B) = \max\{E[e(d)]\pi(d) - k(d), 0\} \quad (30)$$

Using these results, the program evaluates the best action for each set with two eligible ideas remaining, which has the form of equation (28). At each stage, it uses the researcher's beliefs and the revelation procedure to compute the probabilities associated with each state the researcher may find herself in next period.

After working backwards, the program obtains a mapping from every (D, B) state to a best action. With four elements and ten ideas, this program still takes approximately an hour to solve depending on computer processing power. See the appendix for a more detailed description of this program.

7.3. Prizes and Costs

I solve 100 variations of the above problem which differ in researcher beliefs and rewards $\pi(d)$. To obtain the researcher's beliefs about a given pair, I derive α_i and β_i from the following equations, where $\tilde{a}$ is drawn from a beta distribution with $\alpha = 6$ and $\beta = 2$:

$$\begin{aligned}\tilde{a} &= \frac{\alpha_i}{\alpha_i + \beta_i} \\ 0.4 &= \alpha_i + \beta_i\end{aligned}\tag{31}$$

I chose to draw the initial expected affinity from a beta distribution with $\alpha = 6$ and $\beta = 2$ because such a distribution has an expected value of 0.75, but all values above 0.4 occur with relative frequency, as can be seen in Figure 1. This yields a wide array of initial beliefs about pairs.

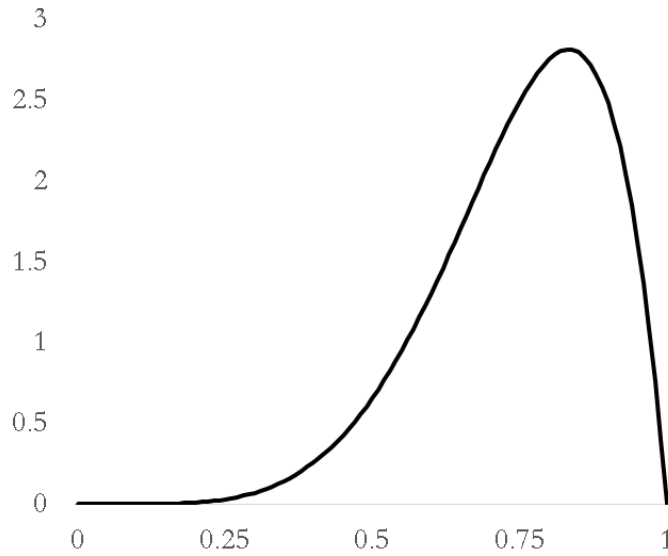


Figure 1: PDF of Beta Distribution with $\alpha = 6$ and $\beta = 2$

I calibrated $\alpha_i + \beta_i = 0.4$ because at this level of certainty, the inherent difficulty of creating more complex ideas can be overcome by learning, at least for a typical case. At this level of certainty, $E[\tilde{a}] = 0.75$, so that $\alpha = 0.3$ and $\beta = 0.1$. Initially, a 3-element idea with $E[a(p)] = 0.75$ for each pair has $E[e(d)] = 0.75^3 \approx 0.42$. Thus, 3-element ideas are initially less likely to succeed than 2-element ideas. However, if the researcher observes one compatibility on each pair, each pair's beta

parameters are increased to $\alpha = 1.3$ and $\beta = 0.1$, so that $E[a(p)] = 1.3 / 1.4 \approx 0.93$. This increases the expected efficacy of the idea to $E[e(d)] \approx 0.93^3 \approx 0.8$, so that such ideas are more likely to succeed than a 2-element idea with $E[a(p)] = 0.75$.

This means a researcher investigating simple ideas composed of a single pair can learn enough to make a 3-element idea just as attractive as a 2-element idea with no information, in the typical case. If the certainty was much higher, learning would not convey much information and Special Case 1 (known affinity) would prevail. Nevertheless, variation in the initial draws of $\tilde{a}$ means I will observe many cases where it is not possible to learn enough to make the efficacy of a three element idea higher than a two-element one (with no information).

Next, I select the value of prizes. There is evidence, both theoretical¹³ and empirical,¹⁴ that important traits about ideas, such as their value, are Pareto distributed. Other studies, however, indicate the distribution of the value of ideas, while fat-tailed and highly skewed, is not Pareto.¹⁵ To address both possibilities, I use two distributions for $\pi(d)$ with the same mean and variance, but where one is a Pareto distribution and the other a log-normal distribution. In practice, I find neither has a meaningful impact on the optimal strategy.

For half the cases, I draw prizes from a Pareto distribution with:

$$\begin{aligned} x_{\min} &= 1 \\ \alpha &= 2.41 \end{aligned} \tag{32}$$

These values imply the median prize has value approximately equal to $4/3$, so that, on average, half of the two-element ideas will satisfy the condition $E[e(d)]\pi(d) > 1$ and therefore be myopically rational to attempt (since the average two-element idea will have $E[e(d)] = 3/4$).

For the other half of the cases, I use a log-normal distribution tuned to have the same median and variance as the Pareto distribution (so that it is primarily the behavior of the tails that differs between the distributions). This implies the log of this distribution follows a normal distribution with $\mu = 0.288$ and $\sigma^2 = 0.633$. Both distributions are plotted in Figure 2.

¹³ Jones (2005) and Kortum (1997) both show many stylized facts about aggregate R&D and growth can emerge if traits of ideas are Pareto distributed.

¹⁴ See Jones (2005), pg. 533-535 for discussion of this evidence.

¹⁵ See Scotchmer (2004), pg. 275-282 for a discussion.

7.4. Simulated Research Decisions

To illustrate features common across many variations, I simulate many actual decisions by a researcher, using the policies that emerge from the solutions to the above problems. To simulate the researcher's problem, I use the researcher's initial beliefs in each model to draw "true" affinities. With the true affinities, I generate the pair compatibilities and efficacies of each idea, as well as the information revealed to the researcher if that idea is attempted. I then have the researcher follow her strategy, observing her choice in each stage. She makes a choice, observes new information, updates her beliefs, and then follows her optimal strategy in the new information and action state. I will perform 1,000 such simulations for each model, and in each simulation the researcher makes 10 choices over 10 periods (after the tenth period she always chooses to quit research).

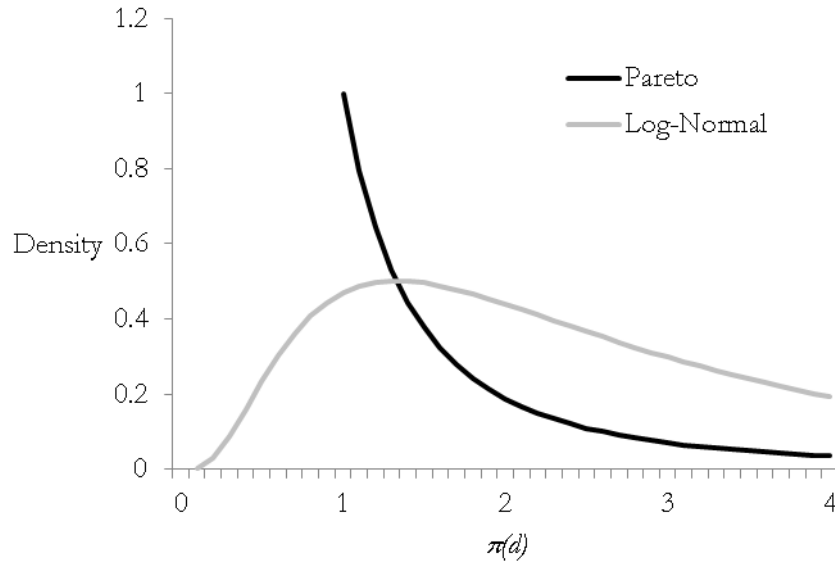


Figure 2: $\pi(d)$ probability density functions

Note: The two pdfs cross again around $\pi(d) = 29$.

8. Numerical Analysis of the General Case

So far, I have made a number of remarks about the characteristics of optimal innovation behavior under simplified settings. In one instance, the learning aspect of the model was suppressed, and in the other, combinatorial features of the model were suppressed. In the general setting, optimal innovation behavior has the characteristic of each.

8.1. Optimal Strategy: Probability A Pair of Elements Are Combined

8.1.1. Functional Forms and Explanatory Variables

I first consider the probability a researcher will attempt to combine two elements as part of a new idea (either a 2-element or 3-element idea). My approach is to model the probability a pair is used in any given decision as a function of characteristics of the pair. In this way, I obtain a profile of the features held by the numerically-derived optimal strategies. Specifically, I run a probit regression with the following form over all the simulated researcher decisions:

$$\Pr(u_{p,t} = 1) = \Phi(\beta_0 + X'_{p,t}\beta) \quad (33)$$

where $u_{p,t}$ is a dummy variable equal to 1 if pair p is used as part of the idea attempted in period t , when the researcher follows an optimal strategy. Each observation corresponds to one pair in one (simulated) researcher's decision. I omit pairs that do not belong to an eligible idea because the probability that $u_{p,t} = 1$ falls to zero in this case. The explanatory variables are traits for each pair, where I choose what traits to include based on the analysis of Remarks 1-7.

Turning first to Section 3, I showed the optimal strategy when affinities are known is straightforward: always choose the idea with the highest expected net value, so long as it exceeds 0 (Remark 1). Therefore, whether or not a pair belongs to the idea with the highest expected net value is a key explanatory variable. I capture this with the dummy variable $\text{Myopic}_{p,t}$, which is equal to 1 if pair $p \in d^*$ and $\arg \max_{d \in D} \{\pi(d)E[e(d)] - k(d)\} = d^*$ in period t .

Of course, when affinities are not known, the optimal strategy in Section 3 is no longer appropriate. Sometimes an idea is selected that is not the myopic best choice, but which provides useful information for future periods. However, when choosing between rival ideas that will provide equally useful information in the future, researchers will still prefer ideas with higher expected net value in the current period, using up the most highly valued ideas first. To account for the fact that the best ideas associated with a pair are fished out over time, I create two variables, $\text{Attempts}_{p,t}$ and $\text{Positive Value}_{p,t}$. The variable $\text{Attempts}_{p,t}$ counts the number of times an idea with pair p has been attempted up to period $t - 1$. As $\text{Attempts}_{p,t}$ rises, there ought to be fewer good ideas left that

contain pair p . The variable $\text{Positive Value}_{p,t}$, conversely, counts the number of remaining eligible ideas that contain pair p and also satisfy $\pi(d)E[e(d)] - k(d) > 0$. The intuition here is that, when the researcher deviates from a myopic strategy, she still wants to minimize her losses. One way to do this is to choose ideas that do not have the highest net expected value, but which still have *high* net expected value. Pairs that belong to many eligible ideas that will be eventually attempted under a pure fishing out strategy are more likely to have *high*, if not the *highest* net expected value.

The optimal strategy for the special case in Section 5 is less straightforward than in the Section 3 case. Researchers adopt a stick-with-the-winner strategy (Remark 7), which I capture by counting the number of times the pair has been observed compatible by the researcher up to period t . I denote this variable $\text{Compatible}_{p,t}$. In Section 5, I also showed that researchers prefer pairs with greater uncertainty (Remark 8). I also capture the degree of certainty about a pair with the variable $\text{Attempts}_{p,t}$ since researchers obtain better information about a pair (usually) when more ideas with it have been attempted.

Unfortunately, a Gittins index strategy does not work in the general setting. Firstly, all ideas are not equally valued, so that the payoff from any one strand of research is declining over time (because of fishing out effects), unless this is offset by the researcher's continual upward assessment of each remaining idea's efficacy. Secondly, in the general setting, all ideas with more than two elements depend on *multiple* pairs. Knowledge that a pair has a high affinity is useless if all other pairs have low affinity, since the researcher can't combine the pair with any others to generate effective ideas. Conversely, if a pair is embedded in a network of many of other pairs with high affinity, learning it too has high affinity is very rewarding, since it can be combined with many other pairs. To measure this effect, I again use the variable $\text{Positive Value}_{p,t}$. Intuitively, any idea with $\pi(d)E[e(d)] - k(d) > 0$ will be tried eventually, and so learning about pairs contained in the idea will have an impact on the value of research. Pairs with a high value of $\text{Positive Value}_{p,t}$ belong to many ideas that would benefit from learning the idea is effective.

The final regression takes the form:

$$\Pr(u_{p,t} = 1) = \Phi(\beta_0 + \beta_1 \cdot \text{Myopic}_{p,t} + \beta_2 \cdot \text{Attempts}_{p,t} + \beta_3 \cdot \text{Positive Value}_{p,t} + \beta_4 \cdot \text{Compatible}_{p,t})$$

(34)

I anticipate the optimal strategy is characterized by the following:

Conjecture 1: Probability of Pairwise Combination. When the probability a researcher will optimally combine two elements as part of a research project is modeled by (34), then $\beta_1, \beta_3, \beta_4 > 0$ and $\beta_2 < 0$.

8.1.2. Results

This is indeed the case, as Table 1 indicates.

Table 1: Probit Regression Characterizing the Optimal Strategy

	Constant	Myopic _{<i>p,t</i>}	Attempts _{<i>p,t</i>}	Positive Value _{<i>p,t</i>}	Compatible _{<i>p,t</i>}
Coefficient	-2.096 (0.002)	2.765 (0.003)	-0.459 (0.008)	0.151 (0.002)	0.428 (0.008)
Observations	3,160,024				
Pseudo R ²	0.626				
Akaike Inf. Crit.	1,183,283				

Note: Standard errors are reported in parentheses.

All the signs are in the anticipated direction, and all parameters are significantly different from zero.

Note that $\text{Positive Value}_{p,t} \geq 1$ whenever $\text{Myopic}_{p,t} = 1$ (otherwise quitting research would be the myopic best choice), that $\text{Positive Value}_{p,t} \leq 3 - \text{Attempts}_{p,t}$ (because the maximum number of eligible ideas associated with a pair is 3), and $\text{Compatible}_{p,t} \leq \text{Attempts}_{p,t}$ (since we can only observe a pair is compatible by attempting an idea containing it). With these constraints, and because these variables are discrete over a small range, I can exhaustively list every feasible combination of pair traits, as well as the probability of selection in Table 2.

Clearly choosing the idea with the highest net expected value is usually the preferred strategy, with the probability of selecting a pair that is the myopic best choice typically on the order of 60-85%. Even when the researcher has twice tried the idea, and never observed a compatibility, such a pair is played 46.1% of the time. Note also that the probability of use is declining in Attempts, although this can be almost perfectly offset by observing compatibilities. For instance, the probability of choosing a pair with Positive Value = 1, Myopic = 0, and Compatible = 0 falls from

2.6% to 0.2% as Attempts rises from 0 to 2, but that it only drops to 2.2% if each attempt reveals the pair to be compatible.

Table 2: Probability of Selection

Myopic	Attempts	Positive Value	Compatibility	$\Pr(u_{x,t} = 1)$
0	0	0	0	0.018
0	0	1	0	0.026
0	0	2	0	0.036
0	0	3	0	0.05
0	1	0	0	0.005
0	1	0	1	0.017
0	1	1	0	0.008
0	1	1	1	0.024
0	1	2	0	0.012
0	1	2	1	0.034
0	2	0	0	0.001
0	2	0	1	0.005
0	2	0	2	0.0155
0	2	1	0	0.002
0	2	1	1	0.007
0	2	1	2	0.022
1	0	1	0	0.794
1	0	2	0	0.834
1	0	3	0	0.869
1	1	1	0	0.641
1	1	1	1	0.785
1	1	2	0	0.696
1	1	2	1	0.827
1	2	1	0	0.461
1	2	1	1	0.629
1	2	1	2	0.776

In the general case, the optimal strategy is a mix between the two strategies discussed in Section 3 and 5. In this model with just four elements, fishing out effects are very strong – at most, any pair only belongs to 3 eligible ideas. Nonetheless, researchers are more likely to return to a pair of elements that has been compatible in the past and they favor learning about pairs that are connected to other good ideas.

8.1.3. Knowledge Accumulation and Fishing Out

These results suggest a natural explanation for how knowledge accumulation effects and fishing out effects can coexist. Within any given set of primitive knowledge elements, the set of ideas that

can be created is finite, and if knowledge is perfect, fishing out effects dominate. This is one reason why the coefficient on the number of ideas attempted using a pair of elements is negative. However, since knowledge is generally not perfect – especially at the outset – knowledge accumulation effects kick in, since learning that elements are compatible tends to expand the set of ideas that can be profitably attempted. This is why the coefficient on the number of compatible observations is positive. Thus, if the set of elements is fixed, knowledge accumulation effects can increase the value of R&D, but only up to an upper bound, given by the perfect knowledge setting. Thereafter, fishing out effects dominate. The only way out of this long-run trap is to expand the set of primitive elements, something this paper does not address (but see Weitzman 1998 for an optimistic take).

8.1.4. Path Dependence

Knowledge accumulation in this model is also related to the concept of technological path dependence. Technological path dependence refers to the idea that certain strands of technology obtain market dominance and hence the attention of future innovators. For example, Acemoglu et al. (2012) present a model where two kinds of technology – carbon neutral and carbon emitting – are substitutes, and in the absence of government policy innovators devote most of their attention to whichever technology has greater market share. Through this mechanism, the transition costs from a carbon emitting to carbon neutral production scheme rise over time, as carbon emitting technology improves at a faster rate than carbon neutral. My model exhibits a similar feature derived from learning, rather than market, effects. Since combinations that have worked in the past are more likely to work in the future, researchers optimally base subsequent research on these combinations, rather than searching for alternatives.

8.1.5. Spillovers

The above model also provides a clear mechanism for how knowledge spillovers might happen. As I have noted above, if the researcher observes some pair p is compatible, this increases the expected efficacy of all other ideas that also include pair p . In this way, positive developments in one idea can spill over to related ideas. There is also a second channel of knowledge spillover though. The probability that pair p forms part of a research project is positively related to the number of ideas containing pair p with $E[e(d)]\pi(d) - k(d) \geq 0$ (captured by the explanatory variable Positive Value). Suppose the researcher observes pair p' is also compatible. This raises the

expected efficacy of all ideas that contain pair p' , including some ideas that *also* contain pair p . If the expected efficacy of such an idea rises by a sufficient amount, it may flip the expected net value of the idea from negative to positive. This increases the probability research projects containing pair p will be attempted. For example, suppose researchers invent a new kind of turbine for power plants. It is known that such a turbine can be reconfigured into a jet engine. This may stimulate research on projects related to jets, but not to turbines at all. In this way, entire technological paradigms can be locked in, since the rise in Positive Value can be temporarily self-reinforcing. Greater knowledge about the affinity of combinations in a subset of elements raises the value of Positive Value for pairs in the subset, which in turn increases the probability of research that increases knowledge about these same pairs.

8.1.6. Radical vs. Incremental Innovation

The above results can also be interpreted in terms of radical versus incremental innovation. The distinction between innovation that generates radically new types of processes and products, and innovation that makes improvements to existing technologies while leaving the basic framework unchanged has roots in economic history.¹⁶ Examples of radical innovation might include the steam engine, electricity, and the computer, while examples of incremental innovation might be a new model of a car or smart phone. The importance of radical innovation for establishing a platform for subsequent improvement is also emphasized in the general purpose technology literature (see Helpman 1998).

In terms of this model, I identify radical innovations as those which are composed of very novel combinations of elements. Such a combination would use pairs with relatively low values of Attempts and Compatibility. In contrast, I identify incremental innovations as those which are composed of relatively common combinations of elements. Such combinations are characterized by pairs with a high degree of certainty about their true affinity. These combinations would use pairs with high values of Attempts and Compatibility.

When radical ideas turn out to be effective, the researcher's beliefs are significantly impacted, and the expected affinity of each pair rises by a comparatively large degree (since beliefs are most responsive to new information when they are characterized by a lot of uncertainty). This is more likely than an incremental innovation to flip some ideas using the same pairs from negative to

¹⁶ See Mokyr (1990), p. 291-292, and Allen (2009), p. 151-155.

positive expected value, and thereby raise Positive Value for pairs. Radical innovation provides one way to temporarily reverse the fishing out of good ideas.

At the same time, it does not follow that a dearth of radical innovation, or “moonshots,” signals a poor outlook for innovation. An assortment of recent works advocate a form of “technological pessimism” (Cowan 2011). This notion is generally based on a raft of arguments, including an intuitive appeal to the reader that technology hasn’t lived up to our dreams. For example, Gordon (2012) asks his readers to consider the impact on their life of losing either (1) all innovation since 2002 or (2) running water and indoor toilets, to demonstrate the paucity of recent innovation. These complaints can be read as frustration with the current incremental state of technological advance, as against the radical innovations that have occurred in the past or which were expected soon.

However, as we can see above, radical innovation (low Attempts and Compatibility) is not generally a superior strategy to incremental innovation (high Attempts and Compatibility). Indeed, it may be worse. As I have discussed above, the coefficients on Attempts and Compatibility approximately offset each other – pairs with no prior attempts are about as likely to be used as those with multiple attempts, so long as each attempt reveals the pair compatible. This approximately equal effect stems from the small number of elements in these simulations, which means each pair can only be used three times. If I increased the number of elements available for combination, the coefficient on Attempts would be reduced relative to the coefficient on compatibility.

Radical innovations, because they rely on untried combinations, are typically riskier and do not benefit much from knowledge accumulation effects. They are always out there, as a backstop research agenda, when other avenues are fished out. Innovations that are very different from the kind around them may be more memorable than another iteration on an existing paradigm, and they may herald promising new avenues of research. But it may be a mistake to complain that there is not more interest in attempting radical innovation. When researchers are disproportionately trying for radical innovations, this is a signal that technological opportunity is low.

8.1.7. Summary

To summarize, researchers overcome the fishing out effect, temporarily, by directing their efforts towards ideas known to succeed. This “stick with the winners” approach is always exhausted in the long run, so that researchers branch out and experiment with combinations that are less studied and which are characterized by uncertainty. If one of these experiments pans out, research along similar lines can continue, as the researcher revises her beliefs about projects she previously

thought to be infeasible. By inducing research on related pairs, the benefits of success can spill over past the initial research project, further sustaining the boom. If none of the experiments pan out though, researchers eventually give up on research and consider the useful ideas which can be pulled from the set of elements exhausted.

8.2. Exploration vs. Exploitation Over Time

The optimal strategy in a no-learning model is myopic and purely exploitative – always choose the idea with the highest net expected value. In contrast, in the learning model, the optimal strategy mixes exploration and exploitation. The researcher is forward looking, so that the possibility of good rewards in the future (conditional on success in the present period) may lead her to attempt ideas that are not myopically optimal.

Over time, a researcher gains information and becomes more confident about the value of each pair's affinity. Concurrently, ideas are fished out and there are fewer ideas to which new information can be applied. Both factors should push researchers towards a purely myopic strategy in later periods, as the value of learning wanes.

This suggests a plausible (but incomplete) analogy between research behavior and a lifecycle model of investment. In a simple investment model, consumption is deferred in early periods to invest (for example, in human capital) until some inflection point when the investor begins to draw down their investments, so that there is nothing left after the terminal period. Just so, it might be presumed, researchers should optimally “invest” in learning in early periods, and then enjoy the fruits of their better information in later periods. This would suggest researchers will start with learning strategies and choose the myopic best choice more often in later periods.

Figure 3 indicates this view is incomplete. To generate this figure, I observe for each simulated researcher decision whether he chooses an exploratory strategy, defined as a choice that does not have the highest expected net value ($E[e(d)]\pi(d) - k(d)$). For each period, I take the fraction of times an exploratory decision is made. I then average these over all 100 variations to generate the chart below. I plot the mean fraction, as well as the 90% upper bound (this gives a boundary below which 90% of variations lie). The unconditional chart includes as observations periods after the researcher has quit research. The other case presents the average fraction of exploratory choices, conditional on the researcher not having quit research in the previous period. Both follow the same general trajectory.

Figure 3 indicates researchers are most likely to make exploratory choices in the *second* period, rather than the first. Some 10% of the time the researcher initially chooses to conduct research on an idea which is not the myopic best choice. This fraction rises to a peak of 16.4% in period 2 before declining to 0% by the final period. The 90% upper bound also follows the same trajectory, rising to a peak in period two and then falling off towards 0% (the 10% lower bound is always 0%).

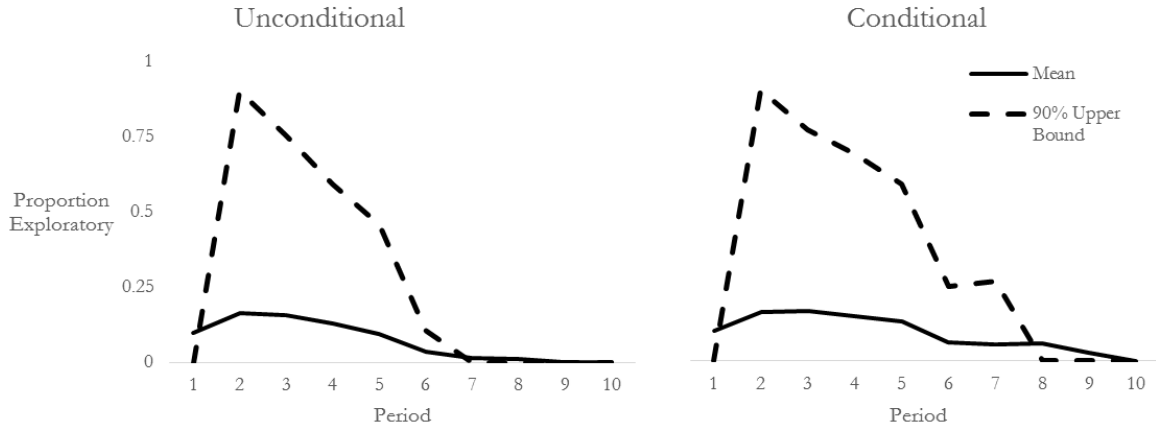


Figure 3: Proportion of Exploratory Choices in Each Period

The conditional chart excludes observations after the researcher has decided to quit research. The unconditional chart does not.

The intuition for this result is as follows. Consider the choice between an exploratory and exploitative strategy in the first period. Choosing the non-exploitative strategy means deferring a higher reward today for the prospect of an even higher reward in the future. The larger the reward offered today, the more unlikely it is the researcher will choose to defer it. And the most valuable ideas are always available in the initial period. Therefore, in the first period, it will tend to be harder for the researcher to adopt a non-exploitative strategy.

In following a myopic strategy, however, the best ideas get used up first. This lowers the penalty of deferring a reward today, and makes an exploratory strategy more attractive. Therefore, after the first few periods, an exploratory strategy is more likely to be adopted.

However, as more research is conducted, an agent becomes more certain about the affinities of different pairs. The importance of learning fades and the optimal strategy should look increasingly like the no-learning strategy again. This suggests the researcher should increasingly favor myopic strategies in later stages.

This discussion is related to the literature on basic versus applied research. While oversimplifying, basic research is associated with projects that expand our knowledge about the

world, but which may be farther from commercial application, while applied research is associated with projects that attempt to commercialize well-understood phenomena. Our model is consistent with a research trajectory that begins applied, proceeds through a relatively long period where basic research is important, and then enters a mature phase where applied research dominates. In our model, we may identify what I have called exploratory research with basic research, since both endure costs in order to learn more (for future application). Myopic research may be then be identified with applied research.

In more informal terms, in the beginning the researcher faces a choice between grabbing low hanging fruit and conducting longer-term research projects. In the beginning, some ideas are so valuable it is worth pursuing them even in the absence of good information. Consider, for example, how the steam engine was developed without a theory of thermodynamics. When these valuable ideas prove successful, they teach the researcher about pair affinities. This implies later researcher may learn much from ideas undertaken for more short-sighted goals (just as physics learned a great deal from steam engines). Later, as these ideas are used up, it becomes increasingly optimal to conduct research before attempting ideas, here represented by the decision to defer the best choice from a myopic point of view in favor of an exploratory strategy. In time though, basic research is no longer necessary.

It is possible to apply this framework to an even larger canvas. Up until the late 19th century, technology proceeded without much guidance from science,¹⁷ because there were enough low-hanging fruit about so that long-term research horizons were not optimal. Petroski (1992), for example, traces the development of a large array of everyday pieces of technology, ranging from tableware to carpentry tools to paperclips. These objects were successively improved by a long line of inventors, each of whose goal in conducting R&D was the improvement of the object at hand, rather than the discovery of knowledge to be applied in distant future contexts.

Smil (2005), however, documents a change in this paradigm during the 19th century. During this period, science increasingly became an input into technology, with electricity, internal combustion engines, chemical industries, and communication infrastructure leading the way. Indeed, Hamilton, Narin and Olivastro (1997) shows that citations to scientific papers by patents have increased significantly over the 1987-1994 period, indicating this trend remains underway in the very recent

¹⁷ See McCloskey (2010), chapter 38 for a discussion relating to the industrial revolution in Britain and Dorn and McClellan III (2006) for a broader perspective.

past. In terms of this model, most of the low-hanging fruit has, at long last, been exhausted, and we have entered the phase where exploratory strategies have become competitive with myopic ones.

The tail-end of this process, if it occurs, is the stuff of science fiction. Vinge (1999), for example, tells us about a future where humanity has colonized the stars, but where genuinely new knowledge is extremely rare. In this science fiction story, what we would consider R&D mostly entails the searching of enormous databases for previous discoveries that can be modified for whatever problem is at hand. This is a world where there is nothing fundamental left to learn, but myopic R&D, drawing on the vast knowledge accumulated over human history, is still practiced.

9. Conclusions

This paper has shown how a model of innovation where researchers learn about the likelihood different pairs of technological “building blocks” work together can generate a number of stylized facts about the innovation process. Knowledge can spill over from one application to another, as researchers observe how the components that comprise each idea interact with each other. This effect also leads to a “standing on the shoulders” of the giants effect, whereby later researchers can develop technologies that would be seen as impractical and unlikely to succeed by earlier researchers. At the same time, for any fixed set of technological building blocks, the set of ideas can be exhausted. Therefore, an optimal strategy mixes exploratory research, whose benefits stem from better information for decision-making in the future, with myopic strategies. The pool of ideas must be periodically restocked by basic research, or else all the fish worth eating will be gone.

Perhaps surprisingly, it is not necessarily optimal to conduct exploratory research immediately, and then use the information gained in later periods. Instead, it is often optimal to initially grab low-hanging fruit that yield a high expected value payoff immediately. Once these are exhausted, basic research becomes a useful strategy. This prediction appears to be in line with the history of technological development writ large.

Going forward, there are several avenues of research that could build on this approach. This paper only discussed a single researcher, who was endowed with knowledge, acting in a partial equilibrium setting. Relaxing these constraints may show this approach can encompass additional facets of the research process. Moreover, this model suggests a way of empirically measuring the state of knowledge for a given technological field, so long as the fields elemental building blocks and the connections between them can be observed. See Clancy 2015 for one such application.

References

- Acemoglu, Daron, Philippe Aghion, Leonardo Bursztyn, and Robert N. Stavins. 2012. "The Environment and Directed Technical Change." *American Economic Review* 102(1) 131-66.
- Arthur, W. Brian. 2009. *The Nature of Technology: What It Is and How It Evolves*. New York: Free Press.
- Auerswald, Philip, Stuart Kauffman, Jose Lobo, and Karl Shell. 2000. "The Production Recipes Approach to Modeling Technological Innovation: An Application to Learning By Doing." *Journal of Economic Dynamics and Control* 24 389-450.
- Bellman, Richard. 1956. "A Problem in the Sequential Design of Experiments." *Sankehyā: The Indian Journal of Statistics (1933-1960)*, Vol. 16, No. 3/4 221-229.
- Borges, Jorge Luis. 1962. *Ficciones (English Translation by Anthony Kerrigan)*. New York, NY: Grove Press.
- Casella, George and Roger L. Berger. 2002. *Statistical Inference: Second Edition*. Pacific Grove, CA: Duxbury Publishing.
- Clancy, Matthew. 2015. *Combinatorial Innovation: Evidence from Two Centuries of Patents*. Dissertation.
- Gittins, John, Kevin Glazebrook, and Richard Weber. 2011. *Multi-Armed Bandit Allocation Indices; 2nd Edition*. Chichester, UK: John Wiley and Sons, Ltd.
- Gittins, John, and You-Gan Wang. 1992. "The Learning Component of Dynamic Allocation Indices." *The Annals of Statistics*, Vol. 20, No. 3 1625-1636.
- Helpman, Elhanan. 1998. *General Purpose Technologies and Economic Growth*. Cambridge, MA: MIT Press.
- Jones, Charles I. 2005. "The Shape of Production Functions and the Direction of Technical Change." *Quarterly Journal of Economics* 120 (2) 517-549.
- Jovanovic, Boyan and Rafael Rob. 1990. "Long Waves and Short Waves: Growth Through Intensive and Extensive Search." *Econometrica* 1391-1409.
- Kauffman, Stuart, Jose Lobo, and William G. Macready. 2000. "Optimal Search on a Technology Landscape." *Journal of Economic Behaviour and Organization* 43 141-166.
- Kortum, Samuel S. 1997. "Research, Patenting and Technological Change." *Econometrica* 1389-1419.
- McClellan, James E. and Harold Dorn. 1999. *Science and Technology in World History: An Introduction*. Baltimore, MD: The John Hopkins University Press.
- McCloskey, Deirdre N. 2010. *Bourgeois Dignity: Why Economics Can't Explain the Modern World*. Chicago, IL: The University of Chicago Press.
- Narin, Francis, Kimberly S. Hamilton, and Dominic Olivastro. 1997. "The Increasing Linkage Between U.S. Technology and Public Science." *Research Policy* 26: 317-330.

- Petroski, Henry. 1992. *The Evolution of Useful Things: How Everyday Artifacts-From Forks and Pins to Paper Clips and Zippers-Came to be as They are*. New York, NY: A. Knopf.
- Poincaré, Henri. 1910. "Mathematical Creation." *The Monist* 321-335.
- Powell, Warren B. 2011. *Approximate Dynamic Programming*. Hoboken, NJ: John Wiley & Sons.
- Ridley, Matt. 2010. *The Rational Optimist*. New York, NY: HarperCollins Publishers.
- Scotchmer, Suzanne. 2004. *Innovation and Incentives*. Cambridge, MA: MIT Press.
- Smil, Vaclav. 2005. *Creating the Twentieth Century*. Oxford, UK: Oxford University Press.
- Vinge, Vernor. 1999. *A Deepness in the Sky*. New York, NY: Tor Publishing.
- Weber, Richard. 1992. "On the Gittins Index for Multiarmed Bandits." *The Annals of Applied Probability* 2(4) 1024-1033.
- Weitzman, Martin L. 1998. "Recombinant Growth." *Quarterly Journal of Economics* 2 331-360.

Appendix A1: Gittins Index

The following proof is based on Weber (1992). Suppose we are in the setting described in Section 5 of this paper.

Optimal strategy with learning: The optimal strategy in every period is to choose the option with the highest Gittins Index $\lambda_i(\alpha_i, \beta_i)$ where

$$\lambda_i(\alpha_i, \beta_i) = \sup\{\lambda_i : v_i(\alpha_i, \beta_i, \lambda_i) = 0\} \quad (35)$$

and $v_i(\alpha_i, \beta_i, \lambda_i)$ given by:

$$\max\left\{\frac{\alpha_i}{\alpha_i + \beta_i}(1 + \delta v_i(\alpha_i + 1, \beta_i, \lambda_i)) + \frac{\beta_i}{\alpha_i + \beta_i}\delta v_i(\alpha_i, \beta_i + 1, \lambda_i) - k - \lambda_i, 0\right\} \quad (36)$$

Proof:

1 – A Single Pair

Suppose pair p_i is the only pair available to choose, so that the researcher's problem collapses to the choice between conducting a research project that includes pair p_i or to quit research. Thus, the researcher's problem can be written as a Bellman equation of the form:

$$v_i(\alpha_i, \beta_i) = \max\left\{\frac{\alpha_i}{\alpha_i + \beta_i}(1 + \delta v(\alpha_i + 1, \beta_i)) + \frac{\beta_i}{\alpha_i + \beta_i}\delta v(\alpha_i, \beta_i + 1) - k, 0\right\} \quad (37)$$

Next, suppose there is an additional charge to conduct research, which I will call the *prevailing charge*, denoted λ_i . The prevailing charge is the same each time the researcher chooses to conduct research on pair p_i . The researcher's problem can then be written as a Bellman equation of the form:

$$v_i(\alpha_i, \beta_i, \lambda_i) = \max\left\{\frac{\alpha_i}{\alpha_i + \beta_i}(1 + \delta v_i(\alpha_i + 1, \beta_i, \lambda_i)) + \frac{\beta_i}{\alpha_i + \beta_i}\delta v_i(\alpha_i, \beta_i + 1, \lambda_i) - k - \lambda_i, 0\right\} \quad (38)$$

Define the *fair charge* $\lambda_i(\alpha_i, \beta_i)$ as the maximum prevailing charge selected so that, in expectation, the optimal strategy makes zero profit:

$$\lambda_i(\alpha_i, \beta_i) = \sup\{\lambda_i : v_i(\alpha_i, \beta_i, \lambda_i) = 0\} \quad (39)$$

Suppose the researcher is facing a fair charge, so that she is indifferent between conducting research and quitting, since both earn expected profit of zero. If the researcher decides to conduct research, she observes the compatibility of pair p_i . If she observe pair p_i to be compatible, her expected profit will be positive going forward, since the charge is fixed but the expected probability of winning a reward in each period is increased. If she observes pair p_i to be incompatible, the prevailing charge will be too high in the next state, so that the researcher would prefer to quit research, earning zero profit in expectation.

Now suppose the prevailing charge is always *lowered* to the fair charge rate, whenever the researcher finds herself in a position where it would be optimal to quit research. This does not affect the researcher's expected profit, since she expects to earn zero under a fair charge, but would have earned zero anyway by quitting research. If the prevailing charge is always reduced in this way, so as to always keep the researcher indifferent when she would otherwise prefer to quit, then the researcher need never stop conducting research. Her expected lifetime profit from such a strategy is zero.

This procedure for reducing the prevailing charge generates a stochastic sequence $\{\lambda_{i,n}\}_{n=0}^{\infty}$ which is nonincreasing in the number of times n the researcher chooses to conduct research.

2 – Many Pairs

Suppose now that there are many pairs available for research, each of which has its own prevailing charge that is periodically reduced in the manner discussed above. Suppose the researcher adopts the following strategy:

Gittins Strategy: In every period, choose the pair with the highest prevailing charge.

The Gittins strategy insures a pair is chosen in every period (since prevailing charges are always lowered when the researcher would otherwise quit research). Such a strategy has zero expected profit. Moreover, there can be *no* strategy that yields strictly positive profit in expectation, since this would require strictly positive profit for at least one pair.

Next, note that the sequences $\{\lambda_{i,n}\}_{n=0}^{\infty}$ associated with each pair are independent of the strategy chosen, since they depend only on the number of times n a pair has been chosen. The Gittins strategy interleaves the many pair sequences into a single nonincreasing sequence of

prevailing charges that maximizes the expected present discounted cost of prevailing charge paid. However, this strategy also yields the maximum expected profit of zero, which means the expected present discounted value of net rewards (absent prevailing charges) must exactly equal the expected cost of charges. Since cost was maximized, this strategy also maximizes rewards, and is therefore an optimal policy.

Since the prevailing charge is periodically lowered to the fair charge, and the fair charge only depends on the state of one pair, an equivalent strategy is to always choose the pair with the highest fair charge, which is given by equation (39).

Appendix A2: Program Details

This section gives some more details on how I solve the general problem presented in Section 7. The program is in the Python language. For clarity of exposition, I will refer to ideas and pairs by the elements $q \in Q$ that ultimately comprise each. To begin, I define the possible sets D in which the researcher may find himself. Since there are 10 ideas, and each idea may be either eligible or ineligible, there are $2^{10} = 1,024$ distinct sets of eligible ideas. For each of these sets, I next define the potential belief vectors of interest. Since I know the initial beta parameters of every $a(p)$, the program can exhaustively list the vectors B that might be attained in any given set.

For example, suppose we are considering the following set

$$D = \{(q_1, q_2, q_3), (q_1, q_2, q_4), (q_1, q_3, q_4), (q_2, q_3, q_4)\} \quad (40)$$

In this set, all of the ideas composed of three elements are eligible, but all of the ideas comprised of two elements are ineligible. In the model, the only way ideas can become ineligible is if they are tried. Therefore, this state can only be arrived at by the researcher after she has conducted research projects on all the two-element ideas. Specifically, we know the researcher has attempted:

$$attempted = \{(q_1, q_2), (q_1, q_3), (q_1, q_4), (q_2, q_3), (q_2, q_4), (q_3, q_4)\} \quad (41)$$

These are the six ideas composed of two elements. Each of these research projects yields information. For each pair, the researcher now has an observation of either one compatibility, or one incompatibility. Therefore, the beta parameters of each pair can take on one of two states: $(\alpha_i + 1, \beta_i)$ or $(\alpha_i, \beta_i + 1)$. Since there are six pairs, and each can take on two states, there are $2^6 = 64$ potential belief vectors associated with the set of eligible ideas.

Often, I can simplify matters by ignoring some pairs. For example, suppose we are considering the following set

$$D = \{(q_1, q_2)\} \quad (42)$$

In this set, only one idea is eligible – all other ideas have already been attempted. This implies a large number of potential belief vectors. For instance, once every idea has been attempted, any given pair can take on 9 states,¹⁸ implying potentially millions of different B vectors. However, most of this information is irrelevant in this case. The only parameters I care about are the ones that describe pairs in eligible ideas. In this case, there is just one pair left in an eligible idea, so I do not have to compute the millions of different B vectors that apply to irrelevant pairs.

Once it has a set of (D, B) states, the program works backwards. It begins by evaluating the null set $D = \{\emptyset\}$, where the only available option is to quit research and earn zero with certainty (for every belief vector B). Next, using this result, the program evaluates the best action for each set with just one eligible idea remaining. When there is just one eligible idea, the problem simplifies to:

$$V(D, B) = \max[E[e(d)]\pi(d) - k(d), 0] \quad (43)$$

Using these results, the program evaluates the best action for each set with two eligible ideas remaining, which has the form of equation (28). At each stage, it uses the researcher's beliefs to compute the probabilities associated with each state the researcher may find herself in next period. For example, suppose the researcher has:

$$\begin{aligned} D &= \{(q_1, q_2), (q_1, q_2, q_3)\} \\ B &= [(0.2, 0.2), (1.2, 0.2), (1.2, 0.2), \dots] \end{aligned} \quad (44)$$

Where the belief parameters apply to pairs (q_1, q_2) , (q_1, q_3) and (q_2, q_3) respectively.

If the researcher chooses research project (q_1, q_2, q_3) , then her possible outcomes are:

¹⁸ For each pair there are two ideas with three elements containing the pair, each of which may reveal compatible, incompatible, or nothing, and one idea with two elements which may reveal compatible or incompatible. Thus, the potential parameter values are:

$(\alpha + 3, \beta), (\alpha + 2, \beta), (\alpha + 2, \beta + 1), (\alpha + 1, \beta), (\alpha + 1, \beta + 1), (\alpha + 1, \beta + 2), (\alpha, \beta + 1), (\alpha, \beta + 2), (\alpha, \beta + 3)$

Table A1: Choosing (q_1, q_2, q_3)

effective?	$V(D, B)$	$B' = B + \omega(d)$	Probability
yes	$\pi(d) + \delta V(D, B') - k(d)$	$[(1.2, 0.2), (2.2, 0.2), (2.2, 0.2), \dots]$	0.37
no	$\delta V(D, B') - k(d)$	$[(0.2, 0.2), (1.2, 0.2), (1.2, 1.2), \dots]$	0.05
no	$\delta V(D, B') - k(d)$	$[(0.2, 0.2), (1.2, 1.2), (1.2, 1.2), \dots]$	0.05
$\vdots$	$\vdots$	$\vdots$	$\vdots$
no	$\delta V(D, B') - k(d)$	$[(1.2, 0.2), (1.2, 1.2), (2.2, 0.2), \dots]$	0.02

In fact, there are 12 potential updated belief vectors that may be attained if the idea is ineffective, reflecting the many different ways an idea can be ineffective (compared to the single way it can be effective). To see where these probabilities come from, consider first the probability the idea is effective, given by the first row of Table A1. Since:

$$E[\Pr(e(d) = 1)] = \prod_{p \in d} E[a(p)] \quad (45)$$

And since

$$E[a(p_i)] = \frac{\alpha_i}{\alpha_i + \beta_i} \quad (46)$$

The probability (q_1, q_2, q_3) is effective is

$$E[e(d)] = \frac{0.2}{0.4} \left(\frac{1.2}{1.4} \right)^2 \approx 0.37 \quad (47)$$

When the idea is effective, each pair is compatible, and so the belief vector next period is given by $[(1.2, 0.2), (2.2, 0.2), (2.2, 0.2), \dots]$, where I have added 1 to the α parameter of each pair in (q_1, q_2, q_3) .

If the idea is ineffective, computing the probability of B' is more involved. Consider the second row, where

$$B' = [(0.2, 0.2), (1.2, 0.2), (1.2, 1.2), \dots] \quad (48)$$

Equation (48) indicates the researcher observed pair (q_2, q_3) to be incompatible, but did not observe any other pairs. The probability of such a revelation requires using the revelation procedure outlined on page 17, and is the joint product of (1) selecting pair (q_2, q_3) , which occurs with $1/3$ probability, and (2) finding the pair is incompatible, which occurs with probability $0.2/1.4$. Since $1/3 \cdot 0.2/1.4 \approx 0.05$, the researcher attaches probability 0.05 to this outcome.

Alternatively, consider the final row, where

$$B' = [(1.2, 0.2), (1.2, 1.2), (2.2, 0.2), \dots] \quad (49)$$

Equation (49) indicates the researcher observed pairs (q_1, q_2) and (q_2, q_3) to be compatible, and pair (q_1, q_3) to be incompatible. There are two ways the revelation procedure could have generated this particular set of observations.

1. It could have (1) drawn pair (q_1, q_2) and found it to be compatible (probability of being drawn is $1/3$, probability of being compatible is $1/2$), (2) drawn pair (q_2, q_3) and found it to be compatible (probability of being drawn is $1/2$ and probability of being compatible is $1.2/1.4$) and (3) drawn pair (q_1, q_3) and found it to be incompatible (probability of being drawn is 1 and probability of being incompatible is $0.2/1.4$). The joint probability of this sequence is approximately 0.01.
2. It could have (1) drawn pair (q_2, q_3) and found it to be compatible (probability of being drawn is $1/3$, probability of being compatible is $1.2/1.4$), (2) drawn pair (q_1, q_2) and found it to be compatible (probability of being drawn is $1/2$, probability of being compatible is $1/2$) and (3) drawn pair (q_1, q_3) and found it to be incompatible (probability of being drawn is 1 and probability of being incompatible is $0.2/1.4$). The joint probability of this sequence is approximately 0.01.

Taken together, the probability of observing this set of observations is approximately 0.02. Similar calculations are performed for each state.

Conversely, if the researcher chooses research project (q_1, q_2) , then her possible outcomes are:

Table A2: Choosing (q_1, q_2)

effective?	$V(D, B)$	$B' = B + \omega(d)$	Probability
yes	$\pi(d) + \delta V(D, B') - k(d)$	$[(1.2, 0.2), (1.2, 0.2), (1.2, 0.2), \dots]$	$\frac{0.2}{0.2 + 0.2} = 0.5$
no	$\delta V(D, B') - k(d)$	$[(0.2, 1.2), (1.2, 0.2), (1.2, 0.2), \dots]$	$\frac{0.2}{0.2 + 0.2} = 0.5$

In this case, because there is just the one pair, the researcher observes either the pair is compatible (with probability $1/2$) or that it is incompatible (also with probability $1/2$).

After working backwards, I have a mapping from every (D, B) state to a best action. With four elements and ten ideas, this program still takes approximately an hour to solve depending on computer processing power.

CHAPTER 4

MANDATES AND THE INCENTIVES FOR ENVIRONMENTAL INNOVATION

Submitted to the *Journal of Environmental Economics and Management*

Matthew S. Clancy and GianCarlo Moschini¹

Abstract

Mandates are policy tools that are becoming increasingly popular to promote renewable energy use. In addition to mitigating the pollution externality of conventional energy, mandates have the potential to promote R&D investments in renewable energy technology. But how well do mandates perform as innovation incentives? To address this question, we develop a partial equilibrium model with endogenous innovation to examine the R&D incentives induced by a mandate, and compare this policy to two benchmark situations: *laissez-faire* and a carbon tax. Innovation is stochastic and the model permits an endogenous number of multiple innovators. We find that mandates can improve upon *laissez-faire*, and that the prospect of innovation is essential for their desirability. However, mandates suffer from several limitations. A mandate creates relatively strong incentives for investment in R&D in low-quality innovations, but relatively weak incentives to invest in high-quality innovations, so that the dispersion of realized innovation quality is comparatively low. Moreover, a mandate achieves lower welfare than a carbon tax, and its optimal level is more sensitive to the structure of the innovation process.

1. Introduction

Given the threat of global climate change resulting from greenhouse gas emissions, there is considerable interest in policies that aim to facilitate the substitution of renewable energy for conventional fossil fuels. In addition to correcting the market failure of a pollution externality, it is recognized that policies that promote adoption of renewable energy also have the potential to affect incentives for innovation in better renewable technologies. Indeed, because of the scale of the problem at hand, the role of research and development (R&D) activities is crucial if a sustainable long-term solution is to be attained (Barrett 2009, Popp 2010). The process of innovation is itself

¹ Matt Clancy and GianCarlo Moschini are equal coauthors.

fraught with market failures, and most policy tools are imperfectly suited to tackle both the pollution mitigation objective and the innovation challenge (Jaffe, Newell and Stavins 2003). A considerable body of work has analyzed the performance of alternative environmental policies vis-à-vis their impact on innovation. The dichotomy of prices versus quantity tools is a recurrent theme in this literature, which has privileged the comparison of carbon taxes with (tradable) pollution permits. The presence of multiple market failures has tended to make ranking of various policies options inconclusive (Fischer, Parry and Pizer 2003), although the balance of evidence reviewed by Requate (2005) appears to favor price-based policies. Existing analytical models, however, do not seem to apply well to a type of quantity tools that has become increasingly popular in recent years: quantity “mandates” that require a certain fraction of consumption to be accounted for by renewable energy.

Mandates set a target for renewable energy production, and it falls upon the producers and suppliers of energy to meet this quota. Renewable portfolio standards are a prominent example of this kind of policy and, as of 2011, were used in 27 US states (Delmas and Montes-Sancho 2011), and six European countries (Haas et al. 2011). Renewable portfolio standards mandate that suppliers of electricity source a set percentage of electricity from renewable sources such as solar, wind, biomass, and hydroelectric providers (Holland 2012). A more direct example is perhaps provided by US biofuel policies. The use of mandates is one of the distinctive features of the 2007 Energy Independence and Security Act (EISA), which envisioned overall biofuel use as transportation fuel in the United States to grow to 36 billion gallons by 2022 (Moschini, Cui and Lapan 2012). Whereas the ability of such mandates to engineer increased adoption of renewable energy is clear, their effectiveness at reducing pollution has been at times controversial, as has the impact on welfare (because of unintended effects, e.g., the food vs. fuel debate) (Janda, Kristoufek and Zilberman 2012). Even less is known about the other critical feature noted above: the ability of corrective policies to induce environmental innovation. In this paper we undertake to study this particular question: just how effective are mandates at promoting innovation?

Interest in the question posed in this paper can be highlighted by the experience with US biofuel policies. Mandates have been effective at spurring the growth of the corn-based ethanol industry, which steadily accumulated the capacity required to produce the mandated targets in a timely fashion. But in order to meet the ambitious targets set out by EISA, a major role is envisioned for advanced biofuels such as cellulosic ethanol: 21 of the 36 billion gallons of biofuels mandated by 2022 are supposed to come from advanced biofuels. The experience, so far, has been disappointing: for several years now, the US Environmental Protection Agency (EPA) has essentially waived the

mandate for cellulosic ethanol. In their latest ruling, for 2014 the EPA proposed to blend a mere 17 million gallons of cellulosic ethanol into the fuel supply, down from the EISA statutory requirement of 1,750 million gallons (EPA 2013). An obvious distinction between corn-based ethanol and cellulosic ethanol is that the former is produced with a mature technology, whereas the latter requires new technological breakthrough to make it scalable and commercially viable. For advanced biofuels, therefore, mandates were really supposed to spur sufficient innovation. Is that a legitimate expectation for a policy tool such as mandates? And, if innovation is the crux of the issue, how do mandates compare with a more standard environmental policy tool such as a carbon tax?

In this paper we analyze the scope of mandates as a tool to promote environmental innovation. Specifically, we consider a market with clean and dirty energy sources that are close substitutes, e.g., renewable energy and fossil fuels. The dirty energy imposes a negative externality on society. The clean energy has no such externality, and the cost of producing it can be lowered through R&D. Following the approach introduced by Parry (1995), Laffont and Tirole (1996) and Denicolo (1999), we view the R&D sector as separate from the production sector adopting the new technology. The profit opportunity that motivates innovators is directly influenced by environmental policies that penalize dirty energy use or reward clean energy use, and it is mediated by patents. The latter are known to permit only imperfect appropriability of the innovation's benefits, which generally leads to under-provision of R&D. Indeed, for environmental innovations where the underlying externality is not fully internalized by private agents, this under-provision problem is believed to be most acute (Jaffe, Newell and Stavins 2005). Because we wish to emphasize the innovation challenges posed specifically by environmental externalities, rather than the general problem of spurring innovation in a market setting, here we take the second-best nature of patents as given and ignore other policy instruments that deal with innovation in general (Clancy and Moschini 2013). In this setting, no environmental policy measure can lead to a first best outcome by itself. The effectiveness of mandates at spurring environmental innovation, therefore, is best understood as compared to a well-defined alternative. Hence, we compare the innovation effects of a mandate with those of a carbon tax (the standard implementation tool of price-based policies).

The incentive to innovate induced by environmental policies has been studied either in a deterministic (Denicolo 1999; Fischer, Parry and Pizer 2003; Scotchmer 2010) or stochastic setting (Biglaiser and Horowitz 1994; Parry 1995; Laffont and Tirole 1996). To fit the distinctive policy challenge of bringing about new technologies such as advanced biofuels, a stochastic framework seems most appropriate. Accordingly, in the model we develop, a firm that invests in R&D gets an

independent draw of a cost-reducing technology for the production of renewable energy. We model both the case of a single innovator and the case of multiple innovators. For the latter case, innovators engage in a form of Bertrand competition, so that the firm with the best innovation is the exclusive licensor, but the price that can be charged is constrained by the firm with the next-best innovation. Multiple innovators can raise welfare through two channels: an increase in the number of innovating firms increases the expected quality of the best innovation that will be discovered, and, the *ex post* royalty rate for the best innovation is reduced by the presence of competitors.² This formulation allows us to capture, in an effective and explicit way, the spillover effect of innovations, and the associated imperfect appropriability problem that is one of the roots of R&D underprovision. Another feature of our model is a plausible presumption about the innovation process: when they choose R&D investments firms have better information than policy makers do when they set the policy. This information asymmetry may stem either from the specialized knowledge of firms, or from the policy-maker's need to set the policy several years in advance, so that it cannot respond to scientific and technological developments (as in the case of the advanced biofuels mandate). To evaluate and compare policies, however, we take the *ex ante* perspective of policy makers who know the distribution of innovation prospects but do not know the actual information possessed by innovators.

Two additional features of our modeling framework deserve a brief discussion. First, we assume that the marginal environmental damage of the externality is constant. This commonly invoked condition, together with the assumption that the conditional distribution of firms' innovation outcomes is uniform, simplifies the analysis considerably and permits the derivation of explicit results. In addition to its analytical attractiveness, this assumption might be appropriate for the case of renewable energy that motivates our analysis. For example, cellulosic ethanol can only address a small portion of the overall energy needs of the economy, and innovations in this area are likely to have a limited impact on the overall level of carbon emission. Furthermore, the energy sector's emissions are small relative to the *cumulative stock* of emissions, which is what ultimately drives climate change. Hence, a linear damage function is arguably appropriate in our context, at least as a local approximation to a convex damage function. A second modeling issue arises because of the

² Because of our focus on R&D incentives, we do not attempt to distinguish between the stages of innovation and diffusion that customarily play distinct roles in this setting (Popp, Newell and Jaffe 2010). In our model diffusion is implicitly assumed to be costless, except for the license fee charged by the innovator).

dynamic implications of R&D incentives. The challenge of devising optimal policies in this context has to deal with Kydland and Prescott's (1977) time consistency problem: once new less-polluting technologies are developed, policy makers might want to change environmental rules, and this *ex post* policy adjustment alters the innovator's *ex ante* incentives. Whether or not policy makers can credibly commit to an environmental policy course, therefore, is of considerable importance (Laffont and Tirole 1996, Denicolo 1999). In our model, having assumed constant marginal environmental damages, the naïve carbon tax that we consider in the analytical section is actually unaffected by the realization of the innovation (Kennedy and Laplante 1999). The optimal mandate and the optimal carbon tax that we consider in the numerical analysis, on the other hand, are not time-consistent. When comparing the performance of mandates with the carbon tax, however, we assume that policymakers have committed to their policy.

Our results show that mandates can in fact improve upon *laissez faire*, and that the prospect of innovation is essential for the desirability of mandates. However, mandates suffer from several limitations and are generally inferior to a carbon tax. We find that a mandate, besides leading to a sub-optimal static equilibrium compared to a carbon tax, also leads to different innovation outcomes. In particular, we find that a mandate is relatively good at incentivizing incremental innovation but a poor spur to breakthrough innovation, as compared with a carbon tax. Mandates also lead to inferior welfare outcomes relative to carbon taxes. In addition to welfare and expected technology results, we highlight the differential distributions of realized technology that different policies can induce. Specifically, we show that carbon taxes induce a more disperse distribution of innovation (either very good or none at all) than a mandate.

The rest of the paper is organized as follows. In Section 2 we lay out our model for the case of a single innovator. Section 3 explicitly compares the mandate policy with the naïve carbon tax when there is one innovator. In Section 4, we extend the model to allow for free entry into the innovation sector. Section 5 compares the mandate policy with the carbon tax when there are multiple innovators. Section 6 uses a numerical simulation to compare the performance of the two policy instruments (mandate and carbon tax) when their level is set "optimally," i.e., accounting for both the externality correction and innovation. This permits us to consider more general welfare conclusions than those derived in the analytical section, and to explore the robustness of our results to the relaxing of certain conditions. We conclude with a summary of our findings, additional discussion of policy implications, and some thoughts about further research.

2. The Model

We model innovation as a purposeful economic activity undertaken by firms seeking to profit from licensing the implementation of their successful ideas. The innovation of interest is modeled as a replacement technology, rather than an abatement technology (another common approach to study environmental innovations). Specifically, we focus on the introduction of a new product that can substitute for an existing product that produces a negative environmental externality. The new product is cleaner—in fact, without much loss of generality, we will assume that this new product has zero emissions. Our modeling of innovation as a replacement technology is consistent with the approach of Denicolo (1999), Laffont and Tirole (1996), and Scotchmer (2011), among others. This approach has also recently been used in the context of climate change (Acemoglu et al. 2012), and fits well the renewable versus conventional fossil fuel context. A distinctive feature of our approach, however, is to model innovation explicitly as a stochastic process. As in Scotchmer (2004), we postulate that innovators decide whether or not to conduct R&D after obtaining a draw from the space of ideas. We extend this approach by assuming that these draws give the innovator only a signal about the *likely* quality of innovation, but the latter remains stochastic. Furthermore, whereas innovators are assumed to observe the signal of technological opportunity prior to making the R&D investment, we presume that the policy setting is determined in advance of the realization of this signal. When comparing alternative policy instruments (carbon tax and mandates), the relevant *ex ante* perspective therefore is that of the policy maker, who knows the distribution of all possible signals but does not observe the realization that drives innovators' decisions. Our stochastic model has other attractive features, including that of permitting an explicit characterization of the multiple innovators setting. Furthermore, this approach is amenable to numerical analysis, which we use to supplement the analytical results.

Consumers are assumed to have quasilinear preferences for a *numeraire* good and energy Q , with the aggregate inverse demand for energy given by $P(Q)$, where $P'(Q) < 0$. There are two forms of energy: an older and dirty form of energy, denoted Q_1 , and a new renewable and clean form of energy, denoted Q_2 . These two sources of energy are perfect substitutes from the consumer's perspective, and thus we can represent total energy used as $Q = Q_1 + Q_2$. Total damage from emission is $X = xQ_1$, where x is the (constant) marginal environmental damage rate. A feature of the innovation context that we wish to model is the fact that the renewable source of energy in question is unlikely to be able to completely supplant the conventional source, and, relative to the

latter, it is expected to be at a scalability disadvantage in both production and distribution. To capture this asymmetry, we assume that the production of the older product displays constant returns to scale at the industry level, whereas renewable energy is produced under decreasing returns to scale at the industry level. Furthermore, whereas the analysis that we present does not restrict the shape of the inverse demand function $P(Q)$, to obtain clear results (especially for the multiple innovators case) we find it convenient to restrict attention to linear industry marginal cost schedules. More specifically, if $C_1(Q_1)$ and $C_2(Q_2, \theta)$ denote the industry cost functions for the two products, conventional energy is assumed to be produced by a perfectly competitive industry with constant marginal cost, i.e.,

$$\frac{\partial C_1(Q_1)}{\partial Q_1} = c_1 \quad (1)$$

whereas the new clean technology displays an upward-sloping marginal cost function:

$$\frac{\partial C_2(Q_2, \theta)}{\partial Q_2} = c_2 - \theta + Q_2 \quad (2)$$

where c_1 and c_2 are fixed parameters, with $c_2 > c_1$, and θ is an index of technological progress that captures the impact of innovation.³ This is illustrated in Figure 1.

³ Linearity of the marginal cost schedule is the main assumption in (2). Conditional on that, setting the slope equal to one is achieved without further loss of generality by choice of the units of measurement for Q .

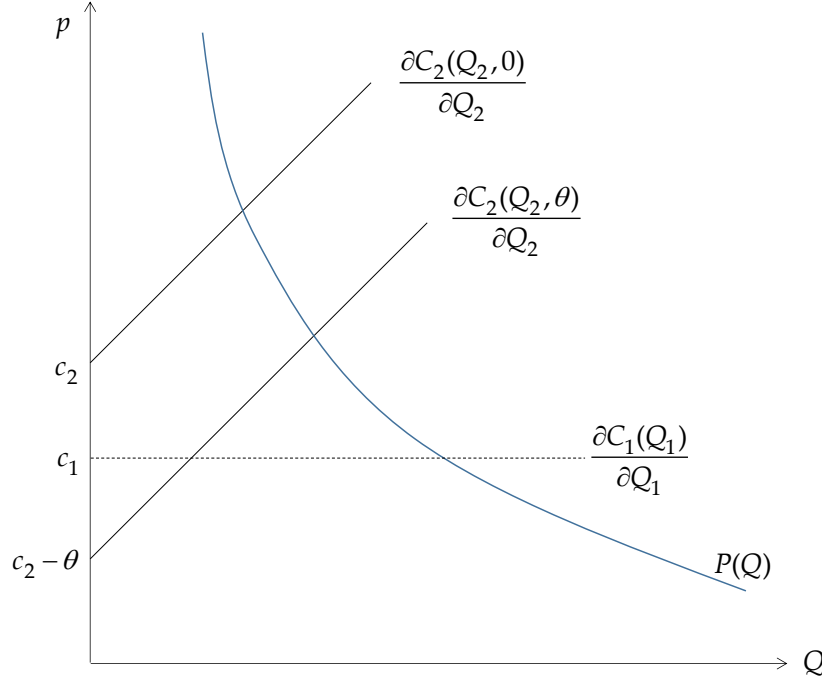


Figure 1. Conventional and renewable energy: Innovation, supply and demand

Initially, $\theta = 0$. The innovation process consist of R&D projects, each of which can produce a draw $\theta \geq 0$ upon incurring a (fixed) cost $k > 0$. As exemplified in (2), innovation lowers the marginal cost of producing renewable energy. More specifically, to develop a new production technology, innovators first receive a draw of ω from a known cumulative probability distribution $G(\omega)$ with domain $[0, \bar{\omega}]$. Given ω , the researcher may choose to pay k to obtain a draw of θ from the conditional distribution function $F(\theta|\omega)$. Whereas the distribution function $G(\omega)$ is unrestricted, apart from the standard monotonicity and continuity properties, the analytical results that we present rely on postulating that $F(\theta|\omega)$ is a uniform distribution. In particular, the density function of this distribution is:

$$f(\theta|\omega) = \begin{cases} 1/\omega & \text{if } \theta \in [0, \omega] \\ 0 & \text{otherwise} \end{cases} \quad (3)$$

The parameter ω characterizes technological opportunity, such that the expected value and upper bound of the innovation draw θ are increasing in ω . But because even the most promising innovation can fail, the lower bound on innovation quality is always zero. We assume that ω is

private information known to innovators, but that the functions $G(\omega)$ and $F(\theta|\omega)$ are common knowledge.

Innovation is understood as producing know-how, and this knowledge is patentable. Innovators produce a blueprint for a new technology, and can license these blueprints to the competitive production sector that produces renewable energy. Licensing is presumed to take the forms of a fixed royalty rate r per unit of Q_2 . In this setting, we want to evaluate the effectiveness of mandates as a policy tool to both ameliorate the externality and promote innovation. For a meaningful benchmark, we compare mandates with a carbon tax and, naturally, both policies with the *laissez faire* situation without any policy. For clarity, we start with the latter cases.

2.1 Innovation under *laissez faire*

For the characterization of both the *laissez faire* situation (absence of government policy), and the case of a carbon tax considered next, the inverse residual demand curve facing clean the production sector can be written as:

$$P_2(Q_2) = \begin{cases} c_1 + t & \text{if } Q_2 \leq P^{-1}(c_1 + t) \\ P(Q) & \text{otherwise} \end{cases} \quad (4)$$

where t denotes the carbon (per unit of dirty energy). For the *laissez faire* case of this section, $t = 0$. In such a case, if clean energy is priced below the cost of dirty energy (c_1), then it captures the entire market; if it is priced above the cost of dirty energy, demand for clean energy falls to zero; and, any quantity $Q_2 \in [0, P^{-1}(c_1)]$ can be sold when clean energy is priced at the cost of dirty energy.

As noted earlier, the realistic scenario is that the new renewable energy source does not completely replace the pre-existing conventional source. That is, the innovation is “nondrastic” in Arrow’s (1962) terminology. The following condition, which we maintain throughout unless otherwise stated, will guarantee this outcome.

Condition 1. The upper bound on technological innovation satisfies $\bar{\omega} \leq c_2 - c_1 + P^{-1}(c_1)$.

In this section we assume that there is only one firm capable of innovating (this assumption is relaxed in later sections). To characterize the innovator's decision problem, consider first the licensing stage for an arbitrary innovation of quality θ . The innovator sets the per-unit royalty r to maximize profits, conditional on the adoption constraint by the competitive clean production sector (which, given the foregoing considerations, faces a perfectly elastic demand at price equal to c_1). Thus, the innovator's optimal royalty maximizes rQ_2 , where the demand from the competitive adopting clean energy sector, for $Q_2 > 0$, satisfies $c_2 - \theta + Q_2 + r = c_1$. When $c_2 - \theta \geq c_1$ there is no strictly positive license fee that can result in any adoption. In such a case, the innovation is insufficient to be cost-competitive with the dirty technology. Thus, licensing only occurs if the innovative step is sufficiently large. More specifically, $\hat{\theta} \equiv c_2 - c_1$ defines the minimum innovative step beyond which the innovation becomes profitable. For $\theta \geq \hat{\theta}$, optimal royalty is $r^* = (\theta - \hat{\theta})/2$, and at this price the quantity licensed is $Q_2 = (\theta - \hat{\theta})/2$. The maximum profit an innovator with technology θ can obtain, when $\theta \geq \hat{\theta}$, is $\pi = (\theta - \hat{\theta})^2/2$ (and, of course, $\pi = 0$ when $\theta < \hat{\theta}$).

A researcher with technological opportunity ω expects the innovation to yield zero profit whenever $\theta < \hat{\theta}$, which happens with probability $\hat{\theta}/\omega$, and thus to make positive profit with probability $1 - \hat{\theta}/\omega$. Expected profit conditional on ω , denoted $\pi(\omega)$, can therefore be written as:

$$\pi(\omega) = \left(1 - \frac{\hat{\theta}}{\omega}\right) \left[\frac{1}{4(\omega - \hat{\theta})} \int_{\hat{\theta}}^{\omega} (\theta - \hat{\theta})^2 d\theta \right] = \frac{(\omega - \hat{\theta})^3}{12\omega} \quad (5)$$

Because the innovator is assumed to be risk neutral, she will choose to conduct research whenever the expected profits from licensing exceed the costs of R&D, which occurs when $\pi(\omega) \geq k$. This implies the existence of a threshold $\hat{\omega} > \hat{\theta}$, which satisfies $\pi(\hat{\omega}) = k$, such that innovation is undertaken if and only if $\omega > \hat{\omega}$.

To understand how innovation affects welfare we note that, given the presumption that innovation is non-drastic, renewable energy is always priced at c_1 (when developed). This means that the total quantity of energy Q , and consumer surplus (denoted S_0), are not affected by innovation. Instead, innovation affects the share of energy produced by renewable sources, and reduces the *status quo ante* damage from externalities (denoted X_0) by xQ_2 . Recall that, when there is

innovation, $Q_2 = (\theta - \hat{\theta})/2$, and thus $E[Q_2] = (\omega - 2\hat{\theta})/4$. License revenues are given in equation (5). Clean producer profits can be shown to be $(\omega - \hat{\theta})^3/24\omega$ in expectation. All told, therefore, expected welfare in the absence of government intervention is

$$E[W] = S_0 - X_0 + \int_{\hat{\omega}}^{\bar{\omega}} \left\{ \left[\frac{(\omega - \hat{\theta})^3}{12\omega} + \frac{(\omega - \hat{\theta})^3}{24\omega} + x \left(\frac{\omega - 2\hat{\theta}}{4} \right) - k \right] \right\} dG(\omega) \quad (6)$$

where the third term in the RHS of (6) is the expected contribution of innovation to welfare.

2.2 The naïve carbon tax

It is well known that *laissez-faire* expected welfare, in this setting, is suboptimal for two reasons: the uncompensated negative externality means there is excess production of dirty fuel, relative to the social optimum; and, the extent of innovation is insufficient. The canonical solution to an externality of the type posited is a carbon tax on the dirty fuel. Because use of fossil fuels incurs a social cost x per unit consumed, if we ignore the prospect of innovation the tax should be set at $t = x$. We will use this “naïve” carbon tax as the benchmark in our analytical results, and consider the optimal carbon tax (which also accounts for the prospect of innovation) in the numerical section. With a unit tax t on fossil fuel, clean producers face an inverse residual demand curve given in (4). As illustrated in the previous section, some innovations may be of insufficient size to be competitive, so that the characterization of the impact of innovation needs to always account for the probability that an innovation of sufficient size actually materializes. To simplify the exposition, and without much loss of generality, it is convenient to maintain the following condition.

Condition 2. The pre-innovation renewable energy technology satisfies $c_2 = c_1 + x$.

This parametric case restricts attention to the situation where the renewable energy source is on the brink of being competitive, given an appropriately tax on the externality posed by the dirty technology. Condition 2 guarantees that the optimal supply of renewable energy is positive for any $\theta > 0$ (i.e., the subsequent analysis can drop the parameter $\hat{\theta}$ corresponding to the minimum inventive step).

Given Condition 2, the optimal license fee and equilibrium quantity of renewable energy satisfy $r^* = Q_2 = \theta/2$. The maximum profit for an innovator possessing an innovation of quality θ , given the existence of the carbon tax t , is: $\pi_t = \theta^2/4$. The innovator's expected profits conditional on technological opportunity, denoted $\pi_t(\omega)$, is given by:

$$\pi_t(\omega) = \frac{\omega^2}{12} \quad (7)$$

Given the existence of a tax t , the threshold $\hat{\omega}_t$ for R&D to be conducted satisfies $\pi_t(\hat{\omega}_t) = k$, and thus $\hat{\omega}_t = \sqrt{12k}$. It is readily verified that this threshold is lower than under *laissez-faire*, i.e., $\hat{\omega}_t < \hat{\omega}$.

Similarly to the *laissez-faire* situation, with a carbon tax a non-drastic innovation does not affect the total quantity of energy nor consumer surplus (here denoted S_0^*). Innovation now improves welfare through its effect on the cost of producing clean fuel, via license profit to the innovator and producer surplus to clean producers. The former was derived in (7). The producer surplus of clean producers can be shown to be $\theta^2/8$, or $\omega^2/24$ in expectation. Combining all elements, expected welfare with innovation, given the carbon tax $t = x$, is:

$$E[W] = S_0^* + \int_{\hat{\omega}_t}^{\bar{\omega}} \left[\frac{\omega^2}{12} + \frac{\omega^2}{24} - k \right] dG(\omega) \quad (8)$$

When compared with (6) we note that the term related to the environmental externality is absent. But welfare is still suboptimal because of the standard appropriability problem that is only partially solved by patents (innovation is underprovided from a social point of view).

2.3 Mandates

A mandate policy specifies a minimum amount of renewable energy to be used as part of the production/consumption portfolio. Such a policy can be modeled either as a proportional or an absolute mandate. With an absolute mandate, distributors must ensure that $Q_2 \geq \hat{Q}$, where $\hat{Q}$ is the mandated minimum quantity of total renewable energy in use. With a proportional mandate, distributors are required to ensure that the total quantity of renewable energy Q_2 amounts to (at

least) a given proportion of the total energy $Q \equiv Q_1 + Q_2$ (e.g., 10 percent). Which of these two modeling approaches are employed does not matter for some questions (e.g., de Gorter and Just 2009, Lapan and Moschini 2012).⁴ In our innovation context, which of these two approaches one uses does entail some modeling differences. Whereas the substantive conclusions one reaches are not affected by this choice, it turns out that an absolute mandate permits a crisper analysis (because it is easier to formalize the results without specifying a particular form for the aggregate energy inverse demand function $P(Q)$). Hence, we proceed by explicitly modeling an absolute mandate.

The implementation of the mandate postulates the existence of a competitive blending sector that combines energy from two sources: conventional energy, sourced at the constant marginal cost c_1 , and renewable energy, priced at its (increasing) marginal cost $\partial C_2 / \partial Q_2 = c_2 - \theta + r + Q_2$. The zero profit condition for the competitive blending sector ensures that, for a given mandate $\hat{Q}$ of renewable energy and corresponding quantity $(Q - \hat{Q})$ of conventional energy, consumers are charged a price $\tilde{P}(Q)$ that is the weighted average of the energy input costs:

$$\tilde{P}(Q) \equiv c_1 \frac{Q - \hat{Q}}{Q} + (c_2 - \theta + r + \hat{Q}) \frac{\hat{Q}}{Q} \quad (9)$$

This formulation presumes that the mandate is binding (typically the case of interest), which is the case whenever the mandate $\hat{Q}$ is such that $(c_2 - \theta + r + \hat{Q}) > c_1$. The issue of feasibility of the mandate should be noted at this juncture. Feasibility is relevant because how much consumers are willing to buy at the blend price is still governed by the (inverse) demand function $P(Q)$. Because consumers (and competitive suppliers) cannot be coerced, not every arbitrary mandate $\hat{Q}$ is feasible. Therefore, in what follows we will assume that the mandate is feasible.

Condition 3. The mandate is feasible in that there is an equilibrium total quantity that solves

$$\tilde{P}(Q^*) = P(Q^*) \text{ and satisfies } Q^* \geq \hat{Q}.$$

⁴ US biofuel mandates can be rationalized as either absolute or proportional mandates. In particular, the EISA legislation specifies absolute mandates for various biofuels. The implementation of such mandates, however, takes the form of proportional mandates imposed on obligated parties. These proportional mandates are set annually by the EPA so that the statutory absolute mandates specified by EISA are met, given the prevailing demand conditions (Schnepf and Yacobucci 2013).

Figure 2 illustrates the case of a feasible mandate ($\hat{Q}'$) and that of an unfeasible mandate ($\hat{Q}''$).

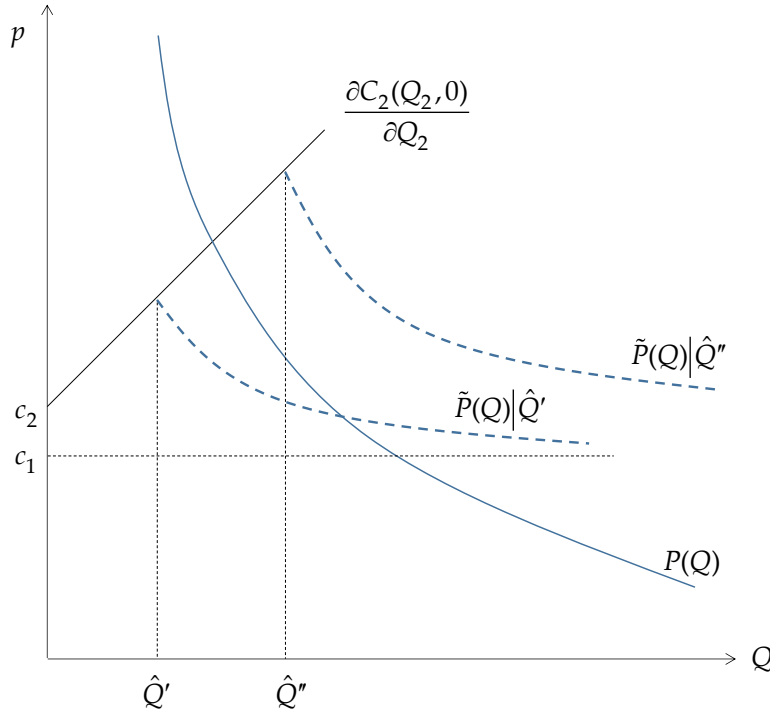


Figure 2. Feasible and unfeasible mandates

If we define $\bar{Q}$ such that $\partial C_2(\bar{Q}, 0)/\partial Q_2 = P(\bar{Q})$, then a sufficient condition to ensure feasibility of the mandate is that $\hat{Q} \leq \bar{Q}$. This requirement is not necessary, however. For a fixed $\hat{Q}$, and given that $c_2 + \hat{Q} > c_1$, the blend price $\tilde{P}(Q)$ is decreasing in Q and asymptotically approaches c_1 from above as Q increases. So, it is quite possible for the equilibrium condition $\tilde{P}(Q^*) = P(Q^*)$ to be satisfied with $Q^* \geq \hat{Q}$ for some $\hat{Q} > \bar{Q}$.⁵

For the analytical results that follow, we wish to concentrate on the policy-relevant case when, post-innovation, the mandate is binding (we will relax this assumption in the numerical analysis of Section 6). In such a case, an innovating firm in possession of technology θ chooses the royalty rate to maximize $r\hat{Q}$, such that $c_2 - \theta + \hat{Q} + r \leq c_2 + \hat{Q}$. This constraint represents the option that clean

⁵ When the mandate exceeds $\bar{Q}$ there may be multiple solutions to $\tilde{P}(Q^*) = P(Q^*)$ (in which case one may appeal to stability conditions to select the relevant equilibrium) or none at all, depending on the shape of the market inverse demand function $P(Q)$.

producers have to meet the mandate by using the pre-innovation technology (for which $\theta = 0$). Clearly, the profit-maximizing license is $r^* = \theta$. The profit attainable by an innovator with technology θ , under a binding mandate, is therefore $\pi_m = \theta \hat{Q}$. What parametric conditions would ensure that the mandate is binding? When the best possible technology is such that $\bar{\omega} \leq \hat{\theta} \equiv c_2 - c_1$, then obviously the mandate is always binding. Otherwise, we note that for the innovator to exceed the mandate the new technology would need to be priced to be competitive with the dirty technology, which is available at marginal cost c_1 . From the *laissez-faire* section 2.1, the profit of the innovator with such a pricing strategy is $\pi = (\theta - \hat{\theta})^2 / 2$. If this profit, for the best possible innovation $\theta = \bar{\omega}$, is no larger than π_m given above, then the mandate will be binding. Hence, when $\bar{\omega} > c_2 - c_1$, to guarantee that the mandate is always binding it suffices to assume:

Condition 4. The mandate satisfies $\hat{Q} \geq (\bar{\omega} - (c_2 - c_1))^2 / 4\bar{\omega}$ and thus is always binding.

Given that the mandate is binding, the expected profit of the innovator with technological opportunity ω , denotes $\pi_m(\omega)$, is:

$$\pi_m(\omega) = \omega \hat{Q} / 2 \quad (10)$$

The lower bound for technological opportunity under which innovation occurs, denoted $\hat{\omega}_m$, solves $\pi_m(\hat{\omega}_m) = k$, and therefore $\hat{\omega}_m = 2k / \hat{Q}$. This threshold is increasing in the cost of R&D and decreasing in the mandate. Under a mandate policy, therefore, R&D occurs with probability $1 - G(\hat{\omega}_m)$. Moreover, by Condition 4, $\hat{\omega}_m \leq \hat{\omega}$, so that the probability of R&D is higher under a mandate than the *laissez-faire* case.

Concerning the impact of innovation on welfare we again find that, because neither the price nor quantity of energy produced changes, there is no change in consumer surplus, now denoted S_0^m , nor the damage from the externality, denoted X_0^m . In this case, the producer surplus Π_0^m of clean firms is also unaffected by innovation, because the innovator fully appropriates the reduction in cost brought about by the innovation. Accordingly, the change in welfare due to innovation is purely derived from licensing profits less R&D costs, so that expected welfare under a mandate is given by:

$$E[W] = S_0^m + \Pi_0^m - X_0^m + \int_{\hat{\omega}_m}^{\bar{\omega}} \left(\frac{\omega \hat{Q}}{2} - k \right) dG(\omega) \quad (11)$$

While the static welfare $S_0^m + \Pi_0^m - X_0^m$ may be suboptimal, the extent of induced innovation, given this policy, is optimal, because the threshold $\hat{\omega}_m$ is exactly determined by where the term under the integral is equal to zero. Moreover, equation (11) implies:

RESULT 1. The welfare maximizing quantity mandate is increased when the social planner takes into account its impact on innovation.

To see why this is the case, consider the case where innovation is not possible. In this case, the optimal static mandate $\hat{Q}_0$ is chosen so that:

$$\frac{\partial S_0^m}{\partial \hat{Q}}_{(-)} + \frac{\partial \Pi_0^m}{\partial \hat{Q}}_{(+)} - \frac{\partial X_0^m}{\partial \hat{Q}}_{(-)} = 0 \quad (12)$$

where the sign of the derivative is given below each term. This first order condition pins down the optimal mandate in the absence of innovation (assuming the usual concavity conditions are satisfied). However, incorporating innovation changes the first order condition such that the optimal mandate, denoted $\hat{Q}_0^I$, now solves:

$$Z(\hat{Q}_0^I) \equiv \frac{\partial S_0^m}{\partial \hat{Q}}_{(-)} + \frac{\partial \Pi_0^m}{\partial \hat{Q}}_{(+)} - \frac{\partial X_0^m}{\partial \hat{Q}}_{(-)} + \int_{\hat{\omega}_m}^{\bar{\omega}} \frac{\omega}{2} dG(\omega)_{(+)} - \left\{ \frac{\omega \hat{Q}}{2} - k \right\}_{(+)} \frac{\partial \hat{\omega}_m}{\partial \hat{Q}}_{(-)} = 0 \quad (13)$$

It is apparent that, when evaluated at $\hat{Q}_0$, $Z(\hat{Q}_0) > 0$. This implies that welfare, when evaluated at $\hat{Q}_0$, is increasing in $\hat{Q}$, so that the mandate should be increased relative to the optimal mandate without innovation. The intuition here is that innovation increases welfare, and a larger mandate increases the incentive to innovate.

3. Mandate vs. Carbon Tax: A Comparison

To evaluate the effectiveness of mandates vis-à-vis the alternative of a carbon tax, it is necessary that the two policies be calibrated to make the comparison meaningful. The benchmark we select initially is to require that the two policies yield the same probability that R&D be undertaken by the innovator, which requires that thresholds of technological opportunity be the same under the two policies, i.e., $\hat{\omega}_t = \hat{\omega}_m$. The threshold under a carbon tax $t = x$ that fully internalizes the social cost of conventional energy, given Condition 2, was shown to be $\hat{\omega}_t = \sqrt{12k}$, whereas under a mandate the innovation threshold is $\hat{\omega}_m = 2k/\hat{Q}$. Hence, the comparable mandate is $\hat{Q} = \sqrt{k/3}$.

Remark 1. When the mandate $\hat{Q}$ is calibrated to yield the same probability of R&D as the carbon tax, the expected value of the technology used by clean producers is the same under both policies.

This observation simply follows from the fact Condition 2 entails that, *ex post*, all technologies are used by clean producers, under either taxes or mandates, for any $\omega \geq \hat{\omega}_t = \hat{\omega}_m$.

Moreover, we also note the following feature of mandates.

Remark 2. As the cost of R&D increases, the level of the mandate must be progressively increased in order to attain the same probability of R&D as a fixed carbon tax.

Whereas the above remarks show that each policy is capable of inducing innovation at the same rate, the welfare implications of these policies differ.

RESULT 2. When a mandate is chosen so that R&D is equally probable under a mandate or a carbon tax, then expected welfare is higher with a carbon tax.

The proof of this result starts by noting that Result 2 will hold so long as:

$$S_0^* + \int_{\hat{\omega}_t}^{\bar{\omega}} \left\{ \frac{\omega^2}{12} + \frac{\omega^2}{24} - k \right\} dG(\omega) \geq S_0^m + \Pi_0^m - X_0^m + \int_{\hat{\omega}_m}^{\bar{\omega}} \left\{ \frac{\omega \hat{Q}}{2} - k \right\} dG(\omega) \quad (14)$$

We note at this juncture that a mandate $\hat{Q} > 0$ is incapable of achieving the first best allocation in the absence of innovation. Given the parameterization of Condition 2, the optimal allocation in the absence of innovation has $P(Q_1) = c_1 + x$ and $Q_2 = 0$. This mix of energy inputs is impossible to achieve with an instrument that only requires $Q_2 \geq \hat{Q}$. Because the first-best allocation (absent innovation) is attained by a carbon tax, it must be that $S_0^* > S_0^m + \Pi_0^m - X_0^m$. Hence, Result 2 will hold when the gains from innovation under the carbon tax exceed those under the mandate, i.e.,

$$\int_{\hat{\omega}_t}^{\bar{\omega}} \left\{ \frac{\omega^2}{12} + \frac{\omega^2}{24} - k \right\} dG(\omega) \geq \int_{\hat{\omega}_t}^{\bar{\omega}} \left\{ \frac{\omega\sqrt{k/3}}{2} - k \right\} dG(\omega) \quad (15)$$

where we have exploited the fact that the mandate is calibrated so that $\hat{\omega}_m = \hat{\omega}_t$, so that the integrals in (14) have the same bounds. A sufficient condition for equation (15) to hold is that the integrand in the LHS exceed the integrand in the RHS for each ω . It is verified that the required condition is

$$\omega/4 \geq \sqrt{k/3} \quad , \quad \forall \omega \in [\hat{\omega}_t, \bar{\omega}] \quad (16)$$

Because this condition is satisfied for the lower bound $\hat{\omega}_t = \sqrt{12k}$, and the LHS in (16) is increasing in ω , the condition is always satisfied.

Given the foregoing, Result 2 continues to hold when the mandate is calibrated so that the probability of R&D is lower than under a carbon tax. However, when mandates are chosen so that the probability of R&D is *higher* than under the carbon tax $t = x$, then it is not apparent whether or not a mandate has lower expected welfare than a carbon tax. In such a case, the gain to welfare from inducing more innovation would need to be weighed against the costs of R&D and distortions to static welfare. We will return to this question in the numerical analysis of Section 6.

4. Multiple innovators

The preceding discussion pertains to the case of a single research firm. In reality, of course, many innovators are typically engaged in competing R&D projects. To model this case, we postulate the existence of a large number of potential innovators, and we assume there is free entry into the renewable energy innovation sector. Innovators are *ex ante* identical and observe a common technological opportunity signal ω . If they choose to conduct R&D, they obtain independent θ

draws from $f(\theta|\omega)$. The innovator who draws the highest θ , denoted θ_1 , has the best technology and becomes the exclusive licensor to the renewable energy production sector. However, the choice of royalty by the innovator who draws θ_1 is now constrained by the presence of competing innovators. Under Bertrand competition, the second-highest θ draw, denoted θ_2 , is the binding constraint. Essentially, as compared with the foregoing analysis, θ_2 plays the same role as the pre-innovation production technique $\theta = 0$ for the single innovator case. But, of course, in the multiple innovator setting θ_2 is endogenous.

To characterize the pricing of innovation with multiple innovators, consider first the *laissez faire* setting. The innovator with the *ex post* best technology θ_1 , presuming that $\theta_1 > \hat{\theta} \equiv c_2 - c_1$, sets the per-unit license r to maximize license profit, conditional on the competitive sector adoption, similar to the single-innovator setting. But here the second best technology θ_2 may limit the price that the licensing innovator can extract. Specifically, the innovator with the best technology maximizes $r(c_1 - c_2 + \theta_1 - r)$, such that $r \leq \theta_1 - \theta_2$. For low realizations of θ_2 , the constraint imposed by the second-best technology does not bind, the single-innovator results continue to hold, and the solution is $r^* = Q_2 = (\theta_1 - \hat{\theta})/2$. Given this unconstrained royalty, it is apparent that the constraint $r \leq \theta_1 - \theta_2$ binds whenever $\theta_2 > (\theta_1 + \hat{\theta})/2$. In such a case the optimal royalty is $r^* = \theta_1 - \theta_2$, and $Q_2 = \theta_2 - \hat{\theta}$. The best innovator's maximum profit, denoted π_1 , is therefore given by:

$$\pi_1 = \frac{(\theta_1 - \hat{\theta})^2}{4} \quad \text{if } \theta_2 \leq (\theta_1 + \hat{\theta})/2 \quad (17)$$

$$\pi_1 = (\theta_1 - \theta_2)(\theta_2 - \hat{\theta}) \quad \text{if } \theta_2 > (\theta_1 + \hat{\theta})/2 \quad (18)$$

The expected profit of a potential entrant now depends on the distribution of θ_1 and θ_2 , which are best described by the concepts of “order statistics” widely used in auction theory (Krishna 2010). Specifically, given n innovators, the probability that an innovator's draw of θ is the maximum draw is equal to the probability that the $n - 1$ other draws are smaller than θ . Because we have assumed a uniform distribution for the innovation projects, then

$$\Pr(\theta = \theta_1 | \omega, n) = (\theta/\omega)^{n-1} \quad (19)$$

Moreover, conditional on a given θ being the maximum draw, the second highest realization θ_2 is the maximum of $n - 1$ independent draws from the uniform distribution on the support of $[0, \theta]$. This is described by the cumulative distribution function:

$$\Pr(\theta_2 < \theta | \omega, n) = (\theta_2 / \theta)^{n-1} \quad (20)$$

which implies that the density function of θ_2 is $((n - 1)/\theta_2)(\theta_2 / \theta)^{n-1}$. Using these results on the distribution of the first and second best innovations, we can determine the expected profitability of participating in the R&D contest. Specifically, with n entrants, the expected profit of each innovator, given technological opportunity ω , can be written as:

$$\pi(\omega, n) = \int_{\hat{\theta}}^{\omega} \left\{ \left(\frac{\theta + \hat{\theta}}{2\theta} \right)^{n-1} \frac{(\theta - \hat{\theta})^2}{4} + \int_{(\theta + \hat{\theta})/2}^{\theta} (\theta - \theta_2)(\theta_2 - \hat{\theta}) \frac{n-1}{\theta_2} \left(\frac{\theta_2}{\theta} \right)^{n-1} d\theta_2 \right\} \left(\frac{\theta}{\omega} \right)^{n-1} \frac{1}{\omega} d\theta \quad (21)$$

This term integrates over the range of values for θ that are both feasible and which earn positive profit. Within the integral, profits are divided into two terms. When $\theta_2 \leq (\theta + \hat{\theta})/2$, which occurs with probability $[(\theta + \hat{\theta})/2\theta]^{n-1}$, profit is given by equation (17). This is the first term under the integral. Conversely, whenever $\theta_2 > (\theta + \hat{\theta})/2$, profit is given by equation (18). This is captured by the second term, itself an integral over possible values of θ_2 .

4.1 Free Entry of Innovators Under a Carbon Tax

With the naïve carbon tax $t = x$, the innovator's problem is similar in structure to the *laissez faire* setting. But here, if the pre-innovation technology is such that Condition 2 applies, it is as if $\hat{\theta} = 0$. Hence, given θ_1 and θ_2 , the best innovator's profit is:

$$\pi_t = \begin{cases} (\theta_1/2)^2 & \text{if } \theta_2 \leq \theta_1 / 2 \\ (\theta_1 - \theta_2)\theta_2 & \text{if } \theta_2 > \theta_1 / 2 \end{cases} \quad (22)$$

Given this conditional profit function, equation (21) can be adapted to yield the expected profit $\pi_t(\omega, n)$ of each innovator facing technological opportunity ω when there are n innovators engaged in R&D:

$$\pi_t(\omega, n) = \int_0^\omega \left\{ \left(\frac{1}{2} \right)^{n-1} \left(\frac{\theta_1}{2} \right)^2 + \int_{\theta_1/2}^{\theta_1} \frac{n-1}{\theta_2} \left(\frac{\theta_2}{\theta_1} \right)^{n-1} (\theta_1 - \theta_2) \theta_2 d\theta_2 \right\} \left(\frac{\theta_1}{\omega} \right)^{n-1} \frac{1}{\omega} d\theta_1 \quad (23)$$

Performing the integration, and simplifying, yields:

$$\pi_t(\omega, n) = \frac{n - (1 - (1/2)^n)}{n(n+1)(n+2)} \omega^2 \quad (24)$$

Note that when $n = 1$ equation (24) reduces to $\omega^2/12$, which is what we found the single innovator profit to be in section 3.2. Profit is clearly increasing in technological opportunity ω . It is also verified that profit is decreasing in the number of innovators n (this occurs for two distinct reasons: as n increases, the probability of any one participant drawing the highest innovations decreases; and, as n increases, the expected royalty for any given innovation decreases).

The equilibrium number of innovators is determined by the free entry condition. In equilibrium, noting that n is an integer, the number of innovators n_t^* satisfies:

$$\pi_t(\omega, n_t^*) \geq k \geq \pi_t(\omega, n_t^* + 1) \quad (25)$$

To emphasize the dependence of the equilibrium number of firms on the R&D outlook parameter ω , which represents the asymmetric information between innovators and policy makers, in what follows this is denoted $n_t^* = n_t(\omega)$.

In section 3.2 we found there existed a unique threshold $\hat{\omega}_t$, where R&D occurred whenever $\omega \geq \hat{\omega}_t$. With free entry, an analogous result can be stated as follows.

Remark 3. Equilibrium with free R&D entry and a carbon tax implies the existence of a sequence of thresholds $\hat{\omega}_t(n)$ such that there are *at least* n active innovators iff $\omega \geq \hat{\omega}_t(n)$.

The threshold levels $\hat{\omega}_t(n)$ are readily computed from (24) and (25):

$$\hat{\omega}_t(n) = \sqrt{\frac{n(n+1)(n+2)}{n-1+(1/2)^n}} k \quad (26)$$

A corollary is that, under free entry, some R&D will take place whenever it is an equilibrium outcome to have at least one innovator, i.e., $\omega \geq \hat{\omega}_t(1)$, which occurs when $\omega \geq \sqrt{12k}$. Naturally, this is the same condition as in the single-innovator case.

As in the single innovator case, consumer surplus is not impacted by innovation, since the price and quantity of final energy is unchanged. Hence, innovation only affects welfare through the profit accruing to the *winning* innovator and the producer surplus of clean producers—denoted $\pi_t^*(\omega, n_t(\omega))$ and $\Pi_t(\omega, n_t(\omega))$, respectively—and total R&D costs $n_t(\omega)k$. Expected welfare can be expressed as:

$$E[W] = S_0^* + \int_0^{\bar{\omega}} \{ \pi_t^*(\omega, n_t(\omega)) + \Pi_t(\omega, n_t(\omega)) - n_t(\omega)k \} dG(\omega) \quad (27)$$

4.2 Free Entry of Innovators Under a Mandate

Under a quota mandate, distributors must ensure $Q_2 \geq \hat{Q}$. The problem faced by an innovator is the same as in section 2.3, but, because of the presumption of Bertrand competition, with the binding constraint now given by the second best technology, where $\theta_2 \geq 0$. This constraint imposes more limitations on the choice of royalty. Specifically, Condition 4 invoked earlier for the single innovator case may no longer suffice to guarantee that the mandate binds. Instead, whenever

$$c_2 - \theta_2 + \hat{Q} < c_1 \quad (28)$$

then the second-best technology is sufficiently good that clean firms would want to use it and exceed the mandate if this technology were competitively priced. The best technology, of course, pre-empts adoption of the second-best technology, but the winning innovator must choose the royalty rate as in the *laissez-faire* setting discussed earlier, in this case leading to an adoption level that exceeds the mandate. The best possible θ_2 is $\bar{\omega}$ and so we assume:

Condition 5: The mandate is large enough to always bind, i.e., $\hat{Q} \geq \bar{\omega} - (c_2 - c_1)$.

This condition is stronger than Condition 4. Note that the best possible θ_1 is also $\bar{\omega}$ and so Condition 5 ensures that the winning innovator cannot choose a royalty such that the mandate is exceeded.

If Condition 5 does not hold, then there will be cases where the winning innovator exceeds the mandate and competes with the fossil fuel alternative as in the *laissez-faire* case. The resulting profits, however, are lower than if fossil fuels were subjected to a carbon tax. Therefore, any draws of θ_1 and θ_2 that lead the winning firm to exceed $\hat{Q}$ are valued less under a mandate than they would be under a carbon tax, and this reduces the incentive to invest in R&D when such draws are possible. Hence, in the analytical derivations that follow we continue to assume that $\hat{Q}$ and $\bar{\omega}$ are such that the mandate is always binding in the post-innovation situation (which provide the best possible set of conditions for the effectiveness of a mandate vis-à-vis the carbon tax). We will return to the possibility that the mandate does not bind in the numerical simulations reported in Section 6.

Given a binding mandate, an innovating firm in possession of the best technology θ_1 chooses the royalty rate to maximize $r\hat{Q}$, such that $c_2 - \theta_1 + \hat{Q} + r \leq c_2 - \theta_2 + \hat{Q}$. The latter indicates the constraint that marginal costs using the best technology, while paying a royalty r , must not exceed marginal costs using the second best technology with a royalty of zero. Clearly, the optimal royalty is $r^* = \theta_1 - \theta_2$ and the quantity induced is $\hat{Q}$. Therefore, the profit of an innovator with the best technology θ_1 , facing the second best technology θ_2 , is:

$$\pi_m = (\theta_1 - \theta_2) \hat{Q} \quad (29)$$

Using the probabilities given by equations (19) and (20), the expected profit of each entrant in the R&D contest, given n innovators and technological opportunity ω , is:

$$\pi_m(\omega, n) = \int_0^\omega \left\{ \int_0^{\theta_1} \frac{n-1}{\theta_2} \left(\frac{\theta_2}{\theta_1} \right)^{n-1} (\theta_1 - \theta_2) \hat{Q} d\theta_2 \right\} \left(\frac{\theta_1}{\omega} \right)^{n-1} \frac{1}{\omega} d\theta_1 \quad (30)$$

After integrating and simplifying, we obtain:

$$\pi_m(\omega, n) = \frac{\hat{Q}}{n(n+1)}\omega \quad (31)$$

Expected profit is increasing in technological opportunity ω and the mandate $\hat{Q}$, and decreasing in the number of innovators. The equilibrium number of innovators $n_m^* = n_m(\omega)$ satisfies:

$$\pi_m(\omega, n_m^*) \geq k \geq \pi_m(\omega, n_m^* + 1) \quad (32)$$

Remark 4. Equilibrium with free R&D entry and a mandate policy implies the existence of a series of thresholds $\hat{\omega}_m(n)$ such that there are *at least* n active innovators iff $\omega \geq \hat{\omega}_m(n)$.

These threshold levels $\hat{\omega}_m(n)$ are computed from (31) and (32):

$$\hat{\omega}_m(n) = \frac{n(n+1)}{\hat{Q}}k \quad (33)$$

Under free entry, there will be some R&D whenever $\omega \geq \hat{\omega}_m(1)$, or $\omega \geq 2k/\hat{Q}$. Naturally, this is the same condition found in section 3.3 for the single innovator case.

Expected welfare under a mandate is no longer straightforward in the presence of multiple innovators. A major difference is that, in the single innovator case, the innovating firm appropriated all of the gains from innovation, so that the price of clean fuel was unchanged by innovation. Under free entry, on the other hand, the winning innovator is only able to appropriate the gains to innovation stemming from improvements over the second best innovation. This also means that the price of clean fuel falls by θ_2 which, because the price consumers pay for energy under a mandate is given by $\tilde{P}(Q)$ in equation (9), leads to an expansion of demand and to a new equilibrium Q' satisfying $\tilde{P}(Q') = P(Q')$, where $Q' > Q^0$ and Q^0 is the pre-innovation equilibrium with mandates. Given a binding mandate, this demand expansion is met entirely by increased dirty fuel production. Whereas the price decline due to the innovation raises consumer surplus, it also increases damages from externalities by $(Q' - Q^0)x$. The increase in damage from the externality exceeds the gain in consumer surplus whenever:

$$(Q' - Q^0)x > \int_{Q^0}^{Q'} P(q) dq - P(Q')Q' + P(Q^0)Q^0 \quad (34)$$

This condition can be rewritten as:

$$\int_{Q^0}^{Q'} [x - (P(q) - P(Q'))] dq > (P(Q^0) - P(Q'))Q^0 \quad (35)$$

The right side of equation (35) is always positive, and so this equation can never be satisfied when $x < (P(Q^0) - P(Q'))$, i.e., when the change in price is greater than marginal damage. Equation (35) is most likely to hold when demand is highly elastic, so that Q' is much larger than Q^0 but the change in price is small relative to marginal damage.

RESULT 3. Under a mandate, the positive welfare impact of innovation is reduced by an expansion of dirty energy consumption. This effect is more sizeable when demand is sufficiently elastic and marginal damage is sufficiently high.

This result establishes that innovation under a mandate is susceptible to a form of the so-called rebound effect. A mandate acts like a tax on the consumption of dirty energy because the use of expensive renewable energy raises the overall cost of energy. But as innovation reduces the cost of renewable energy, the cost of all energy also falls. The increase in total demand is then entirely met by dirty fuel when, as in the case being analyzed, the mandate remains binding.

Under free entry, consumer surplus, clean producer profits, and externalities all depend on the second-best technology, and through this channel their expected values depend on ω . Overall welfare is now written as:

$$E[W] = \int_0^{\bar{\omega}} \left\{ E[S^m | \omega] + E[\Pi^m | \omega] - E[X^m | \omega] + \pi_m^*(\omega, n_m(\omega)) - n_m(\omega)k \right\} dG(\omega) \quad (36)$$

5. Mandate vs. Carbon Tax with Multiple Innovators

With multiple innovators and free entry in the R&D contest, the choice between a carbon tax and a mandate has a greater impact on the character of the realized innovation. To begin, it is more

likely there will be at least n innovators under a carbon tax than under a mandate whenever $\hat{\omega}_m(n) \geq \hat{\omega}_t(n)$. By using equations (26) and (33), and simplifying, this condition reduces to:

$$\frac{k}{\hat{Q}^2} \geq \frac{n+2}{\left(n-1+(1/2)^n\right)n(n+1)} \quad (37)$$

For any given policy the left hand side is fixed, while the right hand side is decreasing in n . This suggests there is a threshold $\hat{n}$ such that at least n innovators are more likely under a carbon tax whenever $n > \hat{n}$, where $\hat{n}$ is defined by:

$$\frac{\hat{n}+1}{\left(\hat{n}-2+(1/2)^{\hat{n}-1}\right)(\hat{n}-1)\hat{n}} \geq \frac{k}{\hat{Q}^2} \geq \frac{\hat{n}+2}{\left(\hat{n}-1+(1/2)^{\hat{n}}\right)\hat{n}(\hat{n}+1)} \quad (38)$$

Because $\hat{\omega}_m(n) \geq \hat{\omega}_t(n)$ for all $n \geq \hat{n}$, and given that $\hat{\omega}_m(n)$ and $\hat{\omega}_t(n)$ are monotonically increasing in n , we conclude with the following result.

RESULT 4. Whenever technological opportunity exceeds a threshold, i.e., $\omega \geq \hat{\omega}_t(\hat{n})$, the number of innovators is (weakly) higher under a carbon tax than under a mandate. Conversely, whenever $\omega \leq \hat{\omega}_t(\hat{n})$, the number of innovators is (weakly) higher under a mandate policy than a carbon tax.

Under either policy, the realized innovation is the best technology drawn by any of the innovators, denoted θ_1 . Conditional on the technology opportunity parameter ω and the number of innovators n , the expected new technology is

$$E[\theta_1 | n, \omega] = \int_0^\omega \theta f_1(\theta | n, \omega) d\theta \quad (39)$$

where $f_1(\theta | n, \omega)$ here is the density function of the distribution of the highest order statistics, which can be related to the primitive distribution $f(\theta | \omega)$ (Krishna 2010). Because of our assumed uniform distribution $f(\theta | \omega) = \theta/\omega$, it follows that

$$f_1(\theta | n, \omega) = n \left(\frac{\theta}{\omega} \right)^{n-1} \frac{1}{\omega} \quad (40)$$

Using this density function and performing the integration in (39) we find:

$$E[\theta_1 | n, \omega] = \frac{n}{n+1} \omega \quad (41)$$

Of course, as discussed in the foregoing, the equilibrium number of innovators will depend on the actual technology opportunity ω and on the policy in place, i.e., $n = n_i(\omega)$, $i = t, m$. Furthermore, from the perspective of a social planner (who does not observe ω), what is relevant is the unconditional expectation of the best technology, that is

$$E[\theta_1] = \int_0^{\bar{\omega}} \frac{n_i(\omega)}{n_i(\omega) + 1} \omega dG(\omega) \quad (42)$$

This makes it apparent that, given the primitive distribution of technological opportunities $G(\omega)$, the expected technology realized depends only on the number of innovators induced by the policy $i = t, m$ for every opportunity ω .

In section 2.4 we showed that, for the single innovator case, setting a mandate equal to $\hat{Q} = \sqrt{k/3}$ ensures that R&D occurs under either policy with equal probability. Using the same policy under free entry preserves this property, but Result 2—according to which the expected technology in use is the same under either policy—is no longer true. When $\hat{Q} = \sqrt{k/3}$, then equation (38) is satisfied by $\hat{n} = 1$ and $n_t(\omega) \geq n_m(\omega)$ for all ω . By equation (42) this implies the expected technology in use will be higher under a carbon tax.

RESULT 5. When the mandate is tuned so that the probability of R&D under a mandate is equal to the probability of R&D under a carbon tax, then the expected technology realized after innovation is better under a carbon tax.

What if the mandate $\hat{Q}$ were tuned so that the expected best technology is the same as under the carbon tax? In order for $E[\theta_1]$ to be the same under either policy, the mandate must be increased from $\sqrt{k/3}$, so that $\hat{\omega}_m(n)$ is decreased. Because $\hat{\omega}_m(n) = n(n+1)k/\hat{Q}$, increasing $\hat{Q}$ will decrease $\hat{\omega}_m(n)$ for all n . Specifically, we will now have $\hat{\omega}_m(1) < \hat{\omega}_t(1)$ so that R&D is more likely to occur under a mandate than under a carbon tax. Moreover, for $E[\theta_1]$ to be the same under either policy, it cannot be that $n_m(\bar{\omega}) > n_t(\bar{\omega})$, where $n_i(\bar{\omega})$ is the number of innovators under

policy i and the best possible technological opportunity. If this were the case, then by Result 4, $n_m(\omega) > n_t(\omega)$ for all $\omega \in [0, \bar{\omega}]$ and by equation (42) $E[\theta_1]$ would be higher under a mandate. Therefore, in this setting, there is some intermediate threshold $\hat{n}$, satisfying $1 < \hat{n} < n_m(\bar{\omega})$, where the number of innovators is higher under a carbon tax for $\omega \geq \hat{\omega}_t(\hat{n})$ and higher under a mandate otherwise. This implies:

RESULT 6. When the mandate is tuned so that the expected best technology is the same under either policy, then the distribution of outcomes under a carbon tax is more disperse than under a mandate.

This result asserts that under a carbon tax there is a higher probability of a very good innovation or none at all. A mandate has a higher probability of *some* innovation, but a lower probability of a very good innovation, since it produces weaker incentives to innovate when technological opportunity is very high.

6. Some Numerical Results

The foregoing analysis has provided some interesting qualitative results on the comparison between mandates, *laissez-faire* and a carbon tax. While these results are illuminating, a limitation is that, apart from Result 2, not much has been said about welfare effects. This is not surprising: as equation (36) indicates, specific welfare conclusions should depend on the particular shape of the demand function $P(Q)$ and on the distribution of technological opportunities $G(\omega)$. Also, our analytical results have been contingent on a few assumptions: that clean energy cannot capture the entire market (Condition 1), that it is on the cusp of being competitive with (taxed) fossil fuels (Condition 2), and that the mandate is always binding (Conditions 4 and 5). In this section we relax these conditions and specify explicit functional forms for $P(Q)$ and $G(\omega)$ so that we may consider welfare effects by means of a numerical analysis.

6.1 Parameterization

We begin by normalizing $c_1 = 100$, so that a tax on dirty energy can be interpreted as a percent of the *laissez-faire* price level. In the baseline parameterization the externality is calibrated to $x = 20$,

so that it amounts to 20% of the private cost of dirty energy,⁶ and we put $c_2 = 120$, consistent with Condition 1 (but this condition does not hold when the marginal damage x is changed from its baseline value). Next, we postulate the inverse demand function $p(Q) = (a - \ln Q)/b$ or, equivalently, that the direct demand function for energy takes the semi-log form:

$$\ln Q = a - bp \quad (43)$$

This is a convenient parameterization which, among other desirable features, can accommodate various hypotheses concerning demand elasticity $\eta \equiv -\partial \ln Q / \partial \ln p$. For this function $\eta = bp$, hence the parameter b can be varied to implement alternative elasticity values. The parameter a is calibrated so that total demand for energy at price $p = c_1$ (and at the baseline elasticity value) is equal to $Q = 100$, that is we put $a = bc_1 + \ln 100$. This normalization means that we can interpret the level of mandates as the percent of total demand under a *laissez-faire* policy. As for $G(\omega)$, we assume that ω is distributed on $[0, \bar{\omega}]$ by an appropriately scaled beta distribution. The probability density function $g(\omega)$ is therefore given by:

$$g(\omega; \alpha, \beta) \propto (\omega / \bar{\omega})^{\alpha-1} (1 - \omega / \bar{\omega})^{\beta-1} \quad (44)$$

where the parameters α and β determine the moments of this distribution and govern its shape. This distribution is very flexible, and alternative choices of α and β can yield both symmetric and skewed density functions. We normalize $\bar{\omega} = 120$ so that, under all possible innovation, the marginal cost of clean energy remains non-negative everywhere.

Given the foregoing functional form assumptions and parametric normalizations, we still have four free parameters that can be varied to gain some insights in the nature of the results. The first of these is the elasticity of demand η . Because this value depends on the evaluation price, for clarity we

⁶This value for the externality cost is meant to be somewhat representative of estimates for the social cost of carbon relative to the cost of transportation fuel. The US government's estimate for the 2015 social cost of carbon, in 2007 dollars, is \$37/ton of CO₂ if a 3% discount rate is used, and \$57/ton of CO₂ if a 2.5% discount rate is used (US Government 2013, p. 3). These discount rates have been criticized for being too high (Johnson and Hope 2012), and so we use the figure associated with the lower 2.5% discount rate as our baseline. Converting this estimate to 2015 dollars yields a social cost of \$65/ton of CO₂. The carbon emission coefficient is 8.9 kg CO₂/gallon of gasoline (EPA 2014), which implies a social cost of carbon is \$0.58 per gallon. Taking the benchmark price of gasoline to be \$3.00/gallon, then the damage imposed by the carbon externality is approximately 20% of the cost of fuel, which is reflected in our baseline value of $x = 20$.

will always measure elasticity with reference to the *laissez-faire* price of energy, where $p = c_1$. For our baseline, we set b so that $\eta = 0.5$. We also consider the cases where $\eta = 0.25$ and $\eta = 1$. Second, we vary the cost of the externality x . As noted, for the baseline we set $x = 20$, but we also consider the cases of $x = 10$ and $x = 40$. Third, we vary the R&D cost k . To calibrate this parameter we relate it to the magnitude of profits that innovation can produce in the *laissez-faire* baseline. Under the highest level of technological opportunity, the expected profit for a single innovator, in view of (5) and the chosen normalizations, is equal to $\pi_t(\bar{\omega}) = 6,250/9$. We consider values of k equal to 3%, 6%, and 12% of this profit level, with 6% corresponding to the baseline.

Fourth, we vary the shape of the distribution of technological opportunity $G(\omega)$. The first moment of the assumed beta distribution is $E[\omega] = \bar{\omega}\alpha/(\alpha + \beta)$. We set $\alpha + \beta = 2$ and, by varying the parameters α and β , we obtain both different values for $E[\omega]$ and different shapes. The baseline parameters are $\alpha = 0.5$ and $\beta = 1.5$, which yield $E[\omega] = 30$. This is a positively skewed distribution (low draws of ω are more likely than high ones), which reflects the belief that incremental innovation is generally more likely than major breakthroughs. The other two cases we consider are $\alpha = 0.25$ and $\beta = 1.75$, which yield $E[\omega] = 15$ (and correspond to an even more positively skewed distribution), and $\alpha = 1$ and $\beta = 1$, which yield $E[\omega] = 60$ (and correspond to a uniform distribution where high draws of ω are equally likely as low ones). These three cases are illustrated in Figure 3.

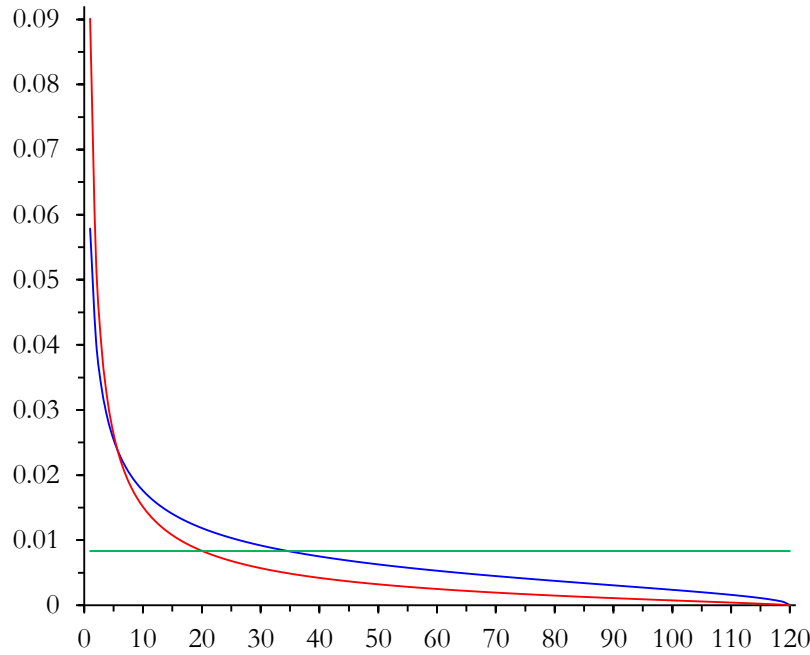


Figure 3. Probability density functions for $G(\omega)$

As for the policies t and $\hat{Q}$, for each set of parameters that we consider, we numerically solve for the value of the policy instrument that maximizes welfare (expected Marshallian surplus).

6.2. Results

The experiments we report, as described in the foregoing, encompass $3^4 = 81$ different parameter combinations. All calculation are coded in Matlab. Some basic descriptive results for the baseline parameters are reported in Table 1. For the single innovator case the expected number of innovators $E[n]$ can be interpreted as the probability that R&D will be conducted. In the baseline setting, under a *laissez-faire* policy, R&D is conducted with probability 0.25 for the single innovator case. The expected quality of innovation $E[\theta_1]$ is 9.66, which improves to 16.05 with multiple innovators. Hence, in either case the “average” technology under *laissez-faire* is insufficient to compete with fossil fuels (the minimum inventive step here is $\hat{\theta} = 20$). Still, some innovation does take place under *laissez-faire*, because some better-than-average draws are viable. The expected quantity of clean energy consumed is small but not negligible, at 2.64 and 8.94 under the single innovator and free entry conditions respectively (recall that the *laissez-faire* quantity of total energy consumed was normalized to 100).

Table 1. Numerical Results for Baseline

	<i>Laissez-Faire</i>		Mandate		Carbon Tax	
	Single Innovator	Free Entry	Single Innovator	Free Entry	Single Innovator	Free Entry
Optimal instrument	-	-	18.65	16.05	23.45	23.40
$E[n]$	0.25	1.52	0.78	2.66	0.56	3.08
$\sqrt{Var(n)}$	0.44	3.10	0.42	2.83	0.50	3.95
$E[\theta_1]$	9.66	16.05	15.63	24.76	14.44	24.17
$\sqrt{Var(\theta_1)}$	20.59	30.15	19.8	28.16	20.45	29.95
$E[Q_2]$	2.64	8.94	18.66	21.68	9.76	23.32
$\sqrt{Var(Q_2)}$	6.99	19.31	0.56	14	9.68	26.98
$E[W]$	126	412	146	455	315	689
$\sqrt{Var(W)}$	421	997	369	993	544	1,114

Note: the baseline parameters are $\eta = 0.5$, $x = 0.2 c_1$, $k = 0.06 \pi(\bar{\omega})$, and $\alpha = 0.5$ and $\beta = 1.5$ (i.e., $E[\omega] = 30$).

An optimal policy (either a mandate or a tax) raises all these quantities, and also improves welfare. The expected quality of innovation $E[\theta_1]$ is also increased significantly, as well as the quantity of clean energy produced. Under an optimal mandate, the probability of R&D more than triples, relative to the *laissez-faire* case, and the expected number of innovators, given free entry, increases from 1.52 to 2.66. Compared with the carbon tax, the mandate induces a greater probability of innovation with a single innovator, but a carbon tax has a higher expected number of entrants when there is free entry. As discussed earlier, this is because a mandate provides comparatively strong incentives to conduct R&D when technological opportunity is low, and this induces firms to enter for more draws of ω than under a carbon tax. The expected profit of R&D increases as ω rises, but it increases at a faster rate for the carbon tax. In the single innovator case, this is irrelevant, since firms make a binary decision to conduct R&D or not. But in the free entry case, the higher profits of a carbon tax can support more innovators, and this leads to a higher overall expected number of entrants (3.08 in a carbon tax, compared to 2.66 under a mandate). The expected quality of innovation, however, is actually higher under a mandate in each case. In the free

entry case, this stems from the differential impact of entrants. Consistent with Result 6, we note that carbon taxes will tend to have more disperse results than the mandate, inducing either many innovators or none at all. Because $\partial E[\theta_1 | \omega, n] / \partial n$ is decreasing in n , the marginal impact of additional entrants under a tax when ω is high (and there are already many firms) is lower than that of additional entrants under a mandate when ω is low (and there are few or no entrants).

The expected quantity of clean energy produced is higher under a mandate, when there is a single innovator, but higher under a carbon tax in the free entry case. However, in both cases, welfare is higher under an optimal carbon tax.⁷ In fact, this is a general numerical result, and we have found it to be true beyond the baseline.

RESULT 7 (NUMERICAL). In all parametric combinations that we considered, expected welfare under the optimal mandate is always lower than under the optimal carbon tax.

Result 7 refers to 81 different parameter combinations, each of which is solved under single innovator and free entry conditions. This result suggests that an optimal mandate, while it improves welfare relative to *laissez-faire*, is inferior to an optimal carbon tax.⁸

To gain further insights into the performance of an optimal mandate, relative to both *laissez-faire* and a carbon tax, Table 2 illustrates the sensitivity of optimal policies to changes in the calibrated parameters.

⁷ Throughout, welfare is measure as expected Marshallian surplus, normalized to zero at the pre-innovation, *laissez-faire* case.

⁸ Indeed, additional simulations, unreported here for space reasons, indicated that even a naïve carbon tax $t = x$ almost always outperforms an optimal mandate in welfare terms.

Table 2. Optimal Policy Instruments Under Alternative Assumptions

	Optimal Mandate			Optimal Carbon Tax		
	No Innovation	Single Innovator	Free Entry	No Innovation	Single Innovator	Free Entry
Baseline	2.4	18.6	16.0	20.0	23.5	23.4
$\eta = 0.25$	1.1	1.5	13.3	20.0	24.4	23.4
$\eta = 1$	5.2	15.2	16.0	20.0	22.5	22.7
$x = 10$	0.0*	0.0*	0.0*	10.0	14.0	14.4
$x = 40$	30.3	41.8	47.0	40.0	47.8	42.8
$k = 0.03\bar{\pi}$	2.4	18.6	16.6	20.0	23.9	22.3
$k = 0.12\bar{\pi}$	2.4	18.1	16.0	20.0	23.0	24.0
$E[\omega] = 15$	2.4	9.5	10.0	20.0	21.7	21.8
$E[\omega] = 60$	2.4	31.4	33.3	20.0	29.3	24.8

Note: Each row changes one parameter, all other parameters as in the baseline.

* reflects rounding (optimal mandates are strictly positive)

The first row of Table 2 reiterates the optimal policies for the baseline parameterization reported in Table 1. Each subsequent row presumes the same parameters as the baseline, except along one dimension. For example, in the second row the parameter that is changed is the elasticity of demand, evaluated at the *laissez-faire* price, which here is $\eta = 0.25$. It is apparent that, across columns, the optimal mandate is considerably more variable than the optimal carbon tax. This suggests the optimal choice of a mandate is sensitive to information about the innovation context, about which policy makers might be less informed than innovators. This conclusion is buttressed by the last two lines of Table 2, which give the optimal policies when the policymaker's beliefs about the outlook for technological innovation are altered. If the outlook for technological opportunity improves from $E[\omega] = 30$ to $E[\omega] = 60$, the optimal mandate increases by 69% and 108% respectively, whereas the optimal carbon tax increases by 25% in the single innovator case and 6% in the free entry case. There is a similarly large divergence when technological opportunity decreases to $E[\omega] = 15$.

Whereas Table 2 illustrates that the magnitude of an optimal policy is more sensitive to information about innovation under a mandate than under a carbon tax, Table 3 illustrates that welfare outcomes are also more sensitive. In this table we decompose the total welfare change $W_1^* - W_0^0$, where W_1^* is the expected welfare with innovation under the optimal instrument choice

(either mandate or carbon tax, all for the multiple-innovators case), and W_0^0 is welfare under *laissez-faire* and no innovation. The decomposition identifies the following four components:

$W_1^0 - W_0^0$: The gain in expected welfare due to innovation under *laissez-faire* ;

$W_0^n - W_0^0$: The “static” gain in expected welfare with a “naïve” level of the policy instrument (i.e., one that does not account for the prospect of innovation);

$(W_1^n - W_0^n) - (W_1^0 - W_0^0)$: The additional gain in expected welfare, relative to *laissez-faire*, due to *policy-supported* innovation (with a naïve level of the instrument);

$W_1^* - W_1^n$: The additional gain in expected welfare from moving to an optimal level of the policy instrument.

Table 3. Welfare Decomposition under Alternative Assumptions (Free Entry)

	Policy	Decomposition				Total
		$W_1^0 - W_0^0$	$W_0^n - W_0^0$	$(W_1^n - W_0^n) - (W_1^0 - W_0^0)$	$W_1^* - W_1^n$	$W_1^* - W_0^0$
Baseline	Tax	412	97	173	7	689
	Mandate	412	2	3	37	455
$\eta = 0.25$	Tax	412	49	173	8	643
	Mandate	412	1	2	25	440
$\eta = 1$	Tax	347	76	98	5	526
	Mandate	347	11	15	35	408
$x = 10$	Tax	323	25	56	9	412
	Mandate	323	0	0	0	323
$x = 40$	Tax	591	575	497	7	1,670
	Mandate	591	400	178	128	1,297
$k = 0.03\bar{\pi}$	Tax	517	97	166	4	784
	Mandate	517	2	5	39	564
$k = 0.12\bar{\pi}$	Tax	288	97	172	9	567
	Mandate	288	2	4	28	322
$E[\omega] = 15$	Tax	179	97	99	3	377
	Mandate	179	2	2	13	196
$E[\omega] = 60$	Tax	1,364	97	317	14	1,791
	Mandate	1,364	2	3	103	1,473

The largest component of the decomposition is the first column, which gives the gain in expected welfare from innovation in a *laissez-faire* setting. This feature is of some interest *per se*, as it emphasizes that the market mechanisms that rationalize the use of policy instruments to spur

innovation also work, to a degree, when no such support is present. Under both policies being considered, the gains from policy-supported innovation, $(W_1^n - W_0^n) - (W_1^0 - W_0^0)$, are generally as large or larger than the gains due to the static increase in allocative efficiency, $W_0^n - W_0^0$. The last column of the decomposition also has direct policy relevance, as it shows how much is gained from moving to an optimal policy that explicitly accounts for the prospect of innovation. This column indicates that the additional gain in welfare from moving to an optimal policy is small under a carbon tax. In the baseline, welfare rises by 97 from the increase in allocative efficiency, by 173 from additional innovation, but only by 7 when moving from the naïve carbon tax of 20 to the optimal carbon tax of 23.4. By contrast, under a mandate, the majority of the gain in welfare stems from moving from the no-innovation policy of 2.4 to the optimal policy of 16.0. A similar result obtains for each of our 9 scenarios: under a mandate, it really is important to tune the policy instrument in response to innovation, whereas with a carbon tax most of the welfare gain can be achieved with the naïve level of the policy instrument.

Table 4 illustrates a final reason why mandates might be unappealing to policymakers if they lack the ability to commit, or they are loss-averse. While expected welfare is improved under all optimal policies, with stochastic innovation it is obviously not true that welfare is improved in all states of the world. When the state of the world is such that the draw θ_1 is low (relative to the signal ω), there may be *ex post* welfare losses. This is illustrated in Table 4, which reports the probability that the realized welfare change relative to the “no-policy, no-innovation” case is negative. For the baseline case, welfare will decrease, as the result of innovation, 76.6% of the time under *laissez-faire*, 42.2% of the times with the mandate, and 0.1% of the times with the carbon tax.

The reasons for a negative welfare change for the *laissez-faire* setting is a quintessential feature of stochastic innovation processes: an R&D investment is a bet that sometimes does not pay off. But the equilibrium nature of the free entry model is such that, when this happens, the size of the loss is limited (for example, in the baseline case the expected welfare change under *laissez-faire*, conditional on being negative, is only -2, whereas the unconditional welfare change in this scenario is 412). These considerations also apply when a policy is in place, but in these settings an additional more powerful effect is at work. That is, to provide enough incentive for innovation, the mandate (and tax) levels need to be appropriately ambitious. In states of the world where the outcome of R&D is poor, however, this means that, *ex post*, the mandate (or tax) is too large. This is particularly a problem when the policy in question is the mandate. As shown in Table 4, the probability of an *ex*

post welfare loss can be quite large under a mandate, whereas it remains small with an optimal carbon tax. Also, the size of the expected welfare loss, conditional on a negative welfare change, is considerably larger with a mandate than with a carbon tax.

Table 4. Undesirable *Ex Post* Welfare Outcomes: Multiple Innovators

Gains from:	Probability of negative welfare change			Excepted welfare change conditional on it being negative, as % of the unconditional expected welfare change		
	<i>Laissez Faire</i>	Optimal Mandate	Optimal Tax	<i>Laissez Faire</i>	Optimal Mandate	Optimal Tax
Baseline	76.6%	42.2%	0.1%	-0.5%	-15.6%	-2.9%
$\eta = 0.25$	76.6%	45.5%	1.1%	-0.5%	-14.1%	-4.0%
$\eta = 1$	76.6%	38.6%	0.4%	-0.6%	-12.5%	-3.6%
$x = 10$	77.1%	77.1%	5.2%	-0.6%	-0.6%	-9.5%
$x = 40$	76.3%	0.0%	0.0%	-0.2%	-1.0%	-1.1%
$k = 0.03\bar{\pi}$	71.3%	37.3%	0.0%	-0.2%	-11.9%	-0.9%
$k = 0.12\bar{\pi}$	82.9%	51.2%	1.2%	-1.0%	-24.8%	-9.0%
$E[\omega] = 15$	88.3%	60.8%	0.1%	-0.6%	-12.2%	-4.5%
$E[\omega] = 60$	3.8%	20.2%	0.2%	-0.3%	-18.5%	-1.2%

Note: Each row changes one parameter, but otherwise maintains the baseline assumptions.

The results of Table 4, which are possible because of the explicit stochastic innovation process that we have characterized, may carry additional policy-relevant implications. On the one hand, as far as choosing between policies, for policymakers it is the *ex ante* perspective that is relevant, in which case the results of Table 4 are immaterial. Yet, the fact that nontrivial *ex post* welfare losses are possible under a mandate suggests that, even for an optimally chosen mandate, there might be strong incentives to change the mandate in such contingencies. That is, the presumption that the government can commit to the policy, which we have maintained in our analysis, might be particularly problematic for the case of mandates. Taking this commitment issue into account in the analysis, of course, would only reinforce our general conclusion, that a mandate policy is less desirable than a carbon tax policy.

7. Conclusion

The direct impact of most environmental policy tools, such as carbon taxes and pollution permits, is to promote the internalization of the external costs of pollution: the reduction of social cost of pollution is achieved by increasing the private cost of (some) agents. It has long been recognized that the privatization of these costs, in addition to ameliorating the externality effect from a static perspective, also has an important dynamic implications because it creates R&D incentives via the so-called induced innovation hypothesis. In this paper we have applied this perspective to the analysis of “mandates,” a policy tool that is becoming increasingly popular in renewable energy contexts.

We find that mandates can in fact improve upon *laissez faire*, and that the prospect of innovation increases the optimal mandate level. The innovation effects are critical and account for most of the desirable welfare impacts of this policy tool. Our numerical results, however, indicate that an optimally calibrated mandate may be much more sensitive to assumptions about innovation, such as the number of potential innovators and the outlook for technological opportunity, than characteristics of demand and the size of the externality. In general, the more promising is innovation, the higher the mandate ought to be. Indeed, the optimal mandate is such that it would typically induce welfare losses in the *status quo* without innovation. In any event, we find that the optimal mandate policy although it is better than *laissez faire*, is clearly dominated by a carbon tax policy.

A novel contribution of our paper, stemming from the explicit stochastic innovation framework that we have developed, is to shed some light on the extent to which alternative policies matter for the distribution of the quality of innovation. In our setting, whereas policies are set based on the entire distribution of possible R&D outlooks, innovators observe a signal on the actual innovation prospects before making their R&D investment. Compared with a mandate, a carbon tax tends to create high profit opportunities when the outlook for R&D turns out to be very good, which induces a flurry of activity that makes the realization of the good innovation outcome likely. Conversely, when the outlook for R&D is weak, mandates may provide more incentive for innovation. Hence, mandates may be a useful policy tool to incentivize R&D when only minor innovations are attainable (or, as for the case of corn-based ethanol mentioned in the introduction, when one deals with a mature technology so that the problem at hand is to promote adoption of existing technologies). But when the goal is to promote breakthrough innovations, as for the cited example of advanced biofuels, it seems that a carbon tax is preferable to mandates. We note that our

qualitative conclusions appear consistent with an emerging empirical literature in renewable energy which shows that quantity-based policies have positive and statistically significant predictors of innovation *only* for older technologies, whereas price-based policies have positive and statistically significant impact for younger technologies (Johnstone, Hascic and Popp 2010).

Whereas mandates may be less effective at spurring innovation for breakthrough technologies, their superior ability to induce innovation when incremental innovation is more likely may make them desirable in some settings. For example, if learning-by-doing is believed to be an important source of technological advance in a field, then it may be more desirable to guarantee that there is *some* kind of innovation, even if it is of low quality, so that the dynamics of learning-by-doing can get started. Alternatively, when innovation proceeds in many incremental steps, mandates may provide higher incentives than a carbon tax for each step in isolation. Whether this translates into a better policy to promote innovation, however, further depends on the extent to which the patent system allows early innovators to capture a sufficient share of the profits from follow-on ideas, a somewhat distinct and complex issue in the economics of intellectual property rights (see, e.g., Scotchmer, 2004, chapter 5).

These caveats notwithstanding, our results indicate that, as a policy to promote environmental innovation, mandates have real limitations. Mandates provide strong incentives for low-quality innovation, but often these are not particularly desirable. Our numerical results further substantiate that welfare with an optimal mandate policy is always lower than with an optimal carbon tax.

References

- Acemoglu, D., P. Aghion, L. Bursztyn, and R.N. Stavins. "The Environment and Directed Technical Change." *American Economic Review* 102(1), 2012: 131-66.
- Barrett, S. "The coming global climate–technology revolution." *The Journal of Economic Perspectives* 23(2)(2009): 53-75.
- Biglaiser, G. and J.K. Horowitz. "Pollution Regulation and Incentives for Pollution-Control Research." *Journal of Economics & Management Strategy* 3.4 (1994): 663-684.
- Clancy, M. and G. Moschini. "Incentives for Innovation: Patents, Prizes, and Research Contracts." *Applied Economic Perspectives and Policy* 35(2)(2013): 206-241.
- de Gorter, H. and D.R. Just. "The Economics of a Blend Mandate for Biofuels." *American Journal of Agricultural Economics*, 91(3), 2009: 738-750.

- Delmas, M.A. and M.J. Montes-Sancho. "US state policies for renewable energy: Context and effectiveness." *Energy Policy* 39(5)(2011): 2273-2288.
- Denicolo, V. "Pollution-reducing innovations under taxes or permits." *Oxford Economic Papers* 51(1), 1999, 184-199.
- EPA. "EPA Proposes 2014 Renewable Fuel Standards, 2015 Biomass-Based Diesel Volume." *Regulatory Announcement*. U.S. Environmental Protection Agency, EPA-420-F-13-048, November 2013.
- EPA. "Greenhouse Gas Emissions from a Typical Passenger Vehicle." *Questions and Answers*. U.S. Environmental Protection Agency, EPA-410-F-14-040a, May 2014.
- Fischer, C., I.W.H. Parry, and W.A. Pizer. "Instrument choice for environmental protection when technological innovation is endogenous." *Journal of Environmental Economics and Management* 45.3 (2003): 523-545.
- Haas, R., G. Resch, C. Panzer, S. Busch, M. Ragwitz, and A. Held. "Efficiency and effectiveness of promotion systems for electricity generation from renewable energy sources—Lessons from EU countries." *Energy* 36(4)(2011): 2186-2193.
- Holland, S.P. "Emissions taxes versus intensity standards: Second-best environmental policies with incomplete regulation." *Journal of Environmental Economics and Management* 63(3)(2012): 375-387.
- Jaffe, A.B., R.G. Newell and R.N. Stavins. "Technological change and the environment." Chapter 11 in: K-G. Mäler and J.R. Vincent, eds. *Handbook of Environmental Economics*. Vol. 1, Elsevier Science, Amsterdam, pp. 461–516.
- Jaffe, A.B., R.G. Newell and R.N. Stavins. "A Tale of Two Market Failures: Technology and Environmental Policy." *Ecological Economics* 54, 2005: 164-174.
- Janda, K., L. Kristoufek and D. Zilberman. "Biofuels: Policies and Impacts." *Agricultural Economics*, 58(8)(2012): 372-386.
- Johnson, Laurie T., and Chris Hope. "The social cost of carbon in U.S. regulatory impact analyses: an introduction and critique." *Journal of Environmental Studies and Science* 2(3) (2012): 205-221.
- Johnstone, Nick, Ivan Hascic, and David Popp. "Renewable Energy Policies and Technological Innovation: Evidence Based on Citation Counts." *Environmental Resource Economics* 45, 2010: 133-155.
- Khishna, V. *Auction Theory*, 2nd edition. Amsterdam: Elsevier, 2010 .
- Kennedy, P.W. and B. Laplante. "Environmental policy and time consistency: emissions taxes and emissions trading." In: E. Petrakis, E.S. Sartzetakis and A. Xepapadeas, eds. *Environmental*

Regulation and Market Power: Competition, Time Consistency and International Trade. Edward Elgar Pub, 1999.

Kydland, F. E. and E.C. Prescott. "Rules rather than discretion: The inconsistency of optimal plans." *Journal of Political Economy* 85(3) (1977): 473-491.

Laffont, J-J. and J. Tirole. "Pollution permits and environmental innovation." *Journal of Public Economics* 62(1), 1996, 127-140.

Lapan, H. and G. Moschini. "Second-Best Biofuel Policies and the Welfare Effects of Quantity Mandates and Subsidies." *Journal of Environmental Economics and Management* 63, 2012: 224-241.

Moschini, G., Cui, J., and Lapan, H. "Economics of Biofuels: An Overview of Policies, Impacts and Prospects," *Bio-based and Applied Economics*, 1(3)(2012): 269-296.

Parry, I.W.H. "Optimal Pollution Taxes and Endogenous Technical Progress." *Resource and Energy Economics* 17, 1995: 69-85.

Popp, D. "Innovation and climate policy." *Annual Review of Resource Economics*, 2 (1)(2010), 275–298.

Popp, D., R.G. Newell and A. Jaffe. "Energy, the environment, and technological change." Chapter 21 in: Hall, B. and N. Rosenberg (Eds.), *Handbook of the Economics of Innovation*, vol. 2 (pp. 873-937). Elsevier, Amsterdam, 2010.

Requate, T. "Dynamic Incentives by Environmental Policy Instruments - A Survey." *Ecological Economics*, 2005: 175-195,

Schnepf, R., and B.D. Yacobucci, "Renewable Fuel Standard (RFS): Overview and Issues," *CRS Report for Congress* 7-5700, Congressional Research Service: Washington, D.C., March 2013.

Scotchmer, S. *Innovation and Incentives*. Cambridge, MA: MIT Press, 2004.

Scotchmer, S. "Cap-and-trade, emissions taxes, and innovation." In: J. Lerner and S. Stern, eds., *Innovation Policy and the Economy*, Volume 11, pp. 29-53. University of Chicago Press, 2011.

US Government. "Technical Support Document: - Technical Update of the Social Cost of Carbon for Regulatory Impact Analysis." *Executive Order 12866*, Interagency Working Group on Social Cost of Carbon, May 2013, Revised November 2013.