

UC Santa Cruz

UC Santa Cruz Previously Published Works

Title

Persistence of false paradigms in low-power sciences

Permalink

<https://escholarship.org/uc/item/3x57s1tf>

Journal

Proceedings of the National Academy of Sciences of the United States of America,
115(52)

ISSN

0027-8424

Authors

Akerlof, George A
Michaillat, Pascal

Publication Date

2018-12-26

DOI

10.1073/pnas.1816454115

Copyright Information

This work is made available under the terms of a Creative Commons Attribution-NonCommercial-NoDerivatives License, available at
<https://creativecommons.org/licenses/by-nc-nd/4.0/>

Peer reviewed



Persistence of false paradigms in low-power sciences

George A. Akerlof^{a,1} and Pascal Michailat^{b,1}

^aMcCourt School of Public Policy, Georgetown University, Washington, DC 20057; and ^bEconomics Department, Brown University, Providence, RI 02912

Contributed by George A. Akerlof, October 31, 2018 (sent for review September 24, 2018; reviewed by Carl Bergstrom and Joshua Graff Zivin)

We develop a model describing how false paradigms may persist, hindering scientific progress. The model features two paradigms, one describing reality better than the other. Tenured scientists display homophily: They favor tenure candidates who adhere to their paradigm. As in statistics, power is the probability (absent any bias) of denying tenure to scientists adhering to the false paradigm. The model shows that because of homophily, when power is low, the false paradigm may prevail. Then, only an increase in power can ignite convergence to the true paradigm. Historical case studies suggest that low power comes either from lack of empirical evidence or from reluctance to base tenure decisions on available evidence.

scientific progress | paradigms | tenure | homophily | power

Friedman (1) argued that while the social sciences lack the experimental evidence available in the physical sciences, this characteristic only slows—but does not prevent—false theories from being “weeded out.” This paper develops a model of science and revisits this argument. We show that when scientists are unbiased, all sciences indeed always converge to the truth, even if they lack empirical evidence to test theories. But there is also homophily, and with homophily in the tenure process, sciences lacking empirical evidence may never converge to the truth. Thus, a lack of experimental evidence may do more than just slow down scientific progress: It may allow worse theories to prevail over better ones.

In modeling science, we follow Kuhn (2). We introduce two paradigms, one giving a better description of the world than the other. Scientists adhere to one or the other paradigm. Scientific inquiry occurs in an environment fashioned after the tenure system in academia: Advisees trained by tenured scientists hope to become tenured scientists themselves. The strength of a paradigm is measured by the fraction of tenured scientists adhering to it; a paradigm is weeded out as more and more tenured scientists adhere to the other paradigm.

The power and significance of a science are defined by using concepts from statistics. Power is the probability of rejecting a worse-paradigm tenure candidate—just as the power of a statistical test is the probability of rejecting a false null hypothesis. Similarly, significance is the probability of rejecting a better-paradigm tenure candidate. How well a science distinguishes between true and false paradigms is measured by the difference between power and significance. Low power is problematic because it reduces this difference, such that false-paradigm scientists are barely more likely to be denied tenure than true-paradigm scientists.

Our central assumption is that scientists have homophilous bias: In the decision to grant tenure, they favor scientists adhering to their paradigm and discriminate against those adhering to the other paradigm. Homophilous bias has been widely documented (3, 4); in fact, even minimal divisions, such as telling groups of boys whether they prefer Klee or Kandinsky, can create strong biases (5). Homophilous bias has been observed in science: People favor others from the same school of thought at every level of academic evaluation—hiring, conference invitations, peer review, tenure evaluations, and awards of grants and honors (6–8). This bias could also explain publication patterns in scientific fields after the early death of a “star”: Collaborators

of the star publish much less, whereas noncollaborators—many of them new entrants to the field—publish much more and are highly cited (9).

Our main result is that a small amount of homophilous bias makes a big difference regarding which paradigm prevails. Without bias, irrespective of power, science always converges to the better paradigm. With homophilous bias, things change: When power is high, science still always converges to the better paradigm, but when power is low, science may converge to the worse paradigm if few scientists initially believe in the better paradigm. Thus, low power does more than just slow down scientific progress: It generates different dynamics, according to which the worse paradigm may prevail. Once a worse paradigm is entrenched, convergence to a better paradigm may only occur if power increases.

That the worse paradigm may prevail at all may be surprising since neither is it more fruitful for research nor do its adherents have greater bias. The worse paradigm may prevail nonetheless because once it has many adherents, a tenure candidate is highly likely to be evaluated by a worse-paradigm scientist. Due to scientists’ homophilous bias, this situation advantages worse-paradigm candidates, to the point that they may become more likely to receive tenure than better-paradigm candidates. Then, the number of tenured scientists believing in the worse paradigm grows faster than the number of tenured scientists believing in the better paradigm: Science moves away from the truth.

Case studies from the history of astronomy, medicine, and economics illustrate the two patterns in our model: entrenchment of inferior paradigms when power is low and convergence toward superior paradigms once power increases—what Kuhn (2) calls scientific revolutions. In these case studies, power is determined both by the data and methods available to test theories and by

Significance

It is believed that a lack of experimental evidence (typical in the social sciences) slows but does not prevent the adoption of true theories. We evaluate this belief using a model of scientific research and promotion in which tenured scientists are slightly biased toward tenure candidates with similar beliefs. We find that when a science lacks evidence to discriminate between theories, or when tenure decisions do not rely on available evidence, true theories may not be adopted. The nonadoption of heliocentric theory in the 16th century, the persistence of bloodletting in the 19th century, the nonadoption of underconsumption theory in the early 20th century, and the persistence of radical mastectomy in the 20th century illustrate such risk.

Author contributions: G.A.A. and P.M. designed research, performed research, and wrote the paper.

Reviewers: C.B., University of Washington; and J.G.Z., University of California, San Diego.

The authors declare no conflict of interest.

This open access article is distributed under [Creative Commons Attribution-NonCommercial-NoDerivatives License 4.0 \(CC BY-NC-ND\)](#).

¹To whom correspondence may be addressed. Email: pascalmichailat@brown.edu or gaa53@georgetown.edu.

This article contains supporting information online at www.pnas.org/lookup/suppl/doi:10.1073/pnas.1816454115/-DCSupplemental.

Published online December 6, 2018.

the norms regarding the appropriate criteria to use in tenure decisions.

Our paper contributes to emerging literature on the theoretical underpinnings of scientific progress. This literature explores how false claims may be accepted as facts if there is a publication bias in favor of positive results (10); how consensus about theories emerges if scientists value conformity (11); how new paradigms are created (12); and how religion influences scientific progress (13).

The population dynamics obtained in this paper are analogous to those that arose when entomologists at the University of Chicago placed populations of two species of flour beetles into jars in the 1950s and 1960s. The entomologists expected that the species more fit to the environment—the species with the higher carrying capacity—would always dominate. But that is not what happened; instead, one or the other species always vanished (14). On close inspection, the reason was determined: Beetles eat their own species' eggs, but are yet more prone to eating the eggs of other species (15). In the flour jars, as in our model, when there is egg-eating bias, two species cannot coexist, and the species doing better in isolation does not necessarily survive in competition.

Model of Science

There are two distinct paradigms: Better and Worse. The Better paradigm gives a more correct description of reality than the Worse, but nobody knows it. Each scientist adheres either to the Better or to the Worse paradigm. At time t , $B(t)$ tenured scientists believe in the Better paradigm, and $W(t)$ believe in the Worse paradigm. The fraction of tenured scientists who believe in the Better paradigm is

$$\sigma(t) = \frac{B(t)}{B(t) + W(t)}. \quad [1]$$

Knowledge is embodied by established scientists, so the strength of a paradigm is measured by the fraction of its adherents among tenured scientists. Since the Better paradigm offers a superior description of the world, knowledge has made progress when $\sigma(t)$ becomes closer to 1.

Scientists' research is based on their respective paradigm: They use their paradigm to explain empirical observations, to guide empirical investigations, and to make theoretical predictions. The quality of a scientist's research is then partially determined by her paradigm. The Better paradigm explains more observations, generates more fruitful empirical investigations, makes more accurate predictions, and can be more easily adjusted to resolve anomalies. Thus, on average, Better scientists produce research of higher quality than Worse scientists. Because no paradigm perfectly describes the real world, however, there is some randomness in research quality. The research process brings additional uncertainty to research quality: Empirical observations are difficult to obtain and subject to measurement error; theoretical explanations and predictions are hard to formulate; and scientists vary in skill, effort, and imagination. So, research quality is noisy and only partially determined by the underlying paradigm.

The initial numbers of scientists adhering to the Better and Worse paradigms are $B(0)$ and $W(0)$. We do not model early adoption, at which stage tenured scientists defect from the existing paradigm and spontaneously adhere to the newly invented paradigm: We take $B(0)$ and $W(0)$ as given. Our focus is on the competition between the Better and Worse paradigms through the tenure system once early adopters start teaching students about their paradigm.

Tenured scientists train advisees at rate $\lambda > 0$. Advisees adhere to the same paradigm as their advisor, and during their

entire career, scientists adhere to the same paradigm. Once an advisee is trained, she becomes an untenured scientist and produces research articulated around her paradigm. Then, she is brought up for tenure. The evaluator of a tenure candidate is randomly chosen from the population of tenured scientists; thus, the candidate is evaluated by a Better scientist with probability $\sigma(t)$ and by a Worse scientist with probability $1 - \sigma(t)$. If she receives tenure, she continues doing research, advises students, and retires at rate δ . If she does not receive tenure, she quits academia.

Let's first consider the case with no homophilous bias in the grant of tenure. Tenure is entirely determined by the quality of research: All candidates whose research quality is above a certain threshold are granted tenure; all others are denied. As research quality is noisy, not all Better candidates are granted tenure and not all Worse candidates are denied. Nevertheless, Better candidates tend to produce higher-quality research than Worse candidates, so they are less likely to be denied tenure. For this reason, we assume that α , the probability of denying tenure to a Better candidate, is lower than $1 - \beta$, the probability of denying tenure to a Worse candidate: $\alpha < 1 - \beta$.

Although the tenure test assesses many attributes of the tenure candidate beside the correctness of her paradigm, the evaluation of a tenure case can be interpreted as a statistical test in which (i) the null hypothesis is that the tenure candidate believes in the more correct paradigm; (ii) the alternative hypothesis is that the tenure candidate believes in the less correct paradigm; (iii) the null hypothesis is rejected when tenure is denied; and (iv) the null hypothesis is accepted when tenure is granted. Then, α is the probability of rejecting the null even though the null is valid, and $1 - \beta$ is the probability of rejecting the null when indeed the null is invalid. Using the analogy with statistics, we introduce two concepts.

Definition 1: The significance of a science is the probability α of denying tenure to a Better scientist when there is no homophilous bias. The power of a science is the probability $1 - \beta$ of denying tenure to a Worse scientist when there is no homophilous bias.

The difference between power and significance measures the gap between the tenure probabilities of Better and Worse scientists when there is no bias; it therefore captures the ability of a scientific field to distinguish between truer and falsar paradigms. The highest difference is achieved if no Worse scientists receive tenure: $1 - \beta = 1$. The lowest difference occurs if all scientists receive tenure with the same probability: $1 - \beta \rightarrow \alpha$. Between these two extremes, a Better scientist is more likely to receive tenure than a Worse scientist, but Worse scientists have some chance of receiving tenure: $\alpha < 1 - \beta < 1$.

Let's now add in the assumption that scientists are biased in favor of those who belong to their paradigm and against those who belong to the alternative paradigm. The agreement or disagreement of belief between the untenured candidate and her tenured evaluator affects tenure decisions: A Better evaluator grants tenure to a Better candidate with higher probability than a Worse evaluator; a Worse evaluator grants tenure to a Worse candidate with higher probability than a Better evaluator. Formally, to the probabilities α and β , we add the homophilous bias $\epsilon \in [0, 1]$. With bias, the tenure decisions are as follows. A Better evaluator denies tenure to a Better scientist with reduced probability $(1 - \epsilon)\alpha$. A Worse evaluator denies tenure to a Better scientist with increased probability $(1 + \epsilon)\alpha$. A Better evaluator grants tenure to a Worse scientist with reduced probability $(1 - \epsilon)\beta$. Finally, a Worse evaluator grants tenure to a Worse scientist with increased probability $(1 + \epsilon)\beta$. To ensure that all probabilities remain in $[0, 1]$, we add the restrictions that $\alpha \leq 1/(1 + \epsilon)$ and $\beta \leq 1/(1 + \epsilon)$.

Tenure probabilities now depend not only on α and β but also on ϵ and σ . With evaluators drawn randomly from the population of tenured scientists, a tenure candidate is evaluated by a Better evaluator with probability σ and by a Worse evaluator with probability $1 - \sigma$. Hence, a Better candidate is denied tenure with probability

$$\hat{\alpha}(\sigma) = \sigma(1 - \epsilon)\alpha + (1 - \sigma)(1 + \epsilon)\alpha = (1 + \epsilon)\alpha - 2\epsilon\alpha\sigma.$$

Similarly, a Worse candidate is granted tenure with probability

$$\hat{\beta}(\sigma) = \sigma(1 - \epsilon)\beta + (1 - \sigma)(1 + \epsilon)\beta = (1 + \epsilon)\beta - 2\epsilon\beta\sigma.$$

Using again the analogy with statistics, $\hat{\alpha}$ is the probability of type I error (false-positive finding) and $\hat{\beta}$ the probability of type II error (false-negative finding).

We measure the difference between the tenure probabilities of Better and Worse scientists with a Youden index:

Definition 2: The Youden index of a science is

$$J(\sigma) = 1 - \hat{\alpha}(\sigma) - \hat{\beta}(\sigma) = 1 - (1 + \epsilon)(\alpha + \beta) + 2\epsilon(\alpha + \beta)\sigma. \quad [2]$$

The Youden index is defined as in statistics: one minus the probability of type I error minus the probability of type II error (16). The Youden index is linearly increasing in $\sigma \in [0, 1]$, from $J(0) = 1 - (1 + \epsilon)(\alpha + \beta)$ to $J(1) = 1 - (1 - \epsilon)(\alpha + \beta)$. It is minimized when $\sigma = 0$ because then all tenured scientists believe in the Worse paradigm, and they are biased against Better tenure candidates and in favor of Worse tenure candidates. It is maximized when $\sigma = 1$ because then all tenured scientists believe in the Better paradigm, and they are biased in favor of Better candidates and against Worse candidates. Thus, if $1 - \beta \geq \alpha + \epsilon/(1 + \epsilon)$, the Youden index is positive for all $\sigma \in (0, 1)$. And if $1 - \beta < \alpha + \epsilon/(1 + \epsilon)$, the Youden index is negative for $\sigma < \sigma^*$, zero at $\sigma = \sigma^*$, and positive for $\sigma > \sigma^*$, where $\sigma^* \in (0, 1/2)$ is defined by

$$\sigma^* = \frac{1}{2} \left[1 - \frac{1 - (\alpha + \beta)}{\alpha + \beta} \cdot \frac{1}{\epsilon} \right]; \quad [3]$$

σ^* will be the key threshold in the dynamics of $\sigma(t)$.

Overall, the model gives a faithful representation of what Kuhn calls revolutionary science—in contrast to normal science (2). Most of the time, scientists engage in normal science: They work within an accepted paradigm, revealed in textbooks and lectures. They use the paradigm to determine important facts; match the paradigm with these facts; and further articulate the paradigm to improve its fit with nature. Our model focuses instead on periods of revolutionary science: when two paradigms compete. Such phases of science arise in response to discovery of anomalies inconsistent with the old paradigm. In these phases, the decision to reject one paradigm is also the decision to accept another.

In our model, scientific knowledge is embodied by established scientists, with the strength of a paradigm indexed by the fraction of scientists adhering to it. This representation accords with Kuhn, who says that a paradigm prevails only after its acceptance by the scientific community (2). Scientific revolutions are battles of old vs. new paradigms for the allegiance of that community.

According to Kuhn, new adherents to a paradigm are not converts from the old one, but are freshly minted scientists. For instance, Kuhn approvingly quotes Planck: “A new scientific truth does not triumph by convincing its opponents and making them see the light, but rather because its opponents eventually

die, and a new generation grows up that is familiar with it” (ref. 2, p. 151). Similarly, in our model, tenured scientists do not switch allegiance, and the ranks of a paradigm grow over time from newly tenured scientists.

Kuhn also emphasizes “resistance” against new paradigms by the adherents of the old paradigm (2). Such resistance is represented in our model by homophilous bias.

Population Dynamics

We now study the dynamics of the population of scientists to determine under which conditions each paradigm—Better or Worse—eventually prevails.

The population dynamics are the outcome of a horse race regarding which type of scientist is growing at the faster rate. Indeed, differentiating [1] with respect to t , we find that the fraction of Better scientists in the population of tenured scientists evolves according to

$$\dot{\sigma}(t) = \sigma(t) [1 - \sigma(t)] [g^B(t) - g^W(t)], \quad [4]$$

where $g^B(t) = \dot{B}(t)/B(t)$ is the growth rate of Better tenured scientists, and $g^W(t) = \dot{W}(t)/W(t)$ is the growth rate of Worse tenured scientists. The growth rate $g^B(t)$ is simply

$$g^B(t) = \lambda [1 - \hat{\alpha}(\sigma(t))] - \delta.$$

The first term reflects that Better scientists train advisees at rate λ , and with probability $1 - \hat{\alpha}(\sigma(t))$ the advisees are granted tenure. The second term reflects that tenured scientists retire at rate δ . Similarly,

$$g^W(t) = \lambda \hat{\beta}(\sigma(t)) - \delta.$$

The first term reflects that Worse scientists train advisees at rate λ , and with probability $\hat{\beta}(\sigma(t))$ such advisees are granted tenure. The second term reflects retirement. Combining these equations yields

$$\dot{\sigma}(t) = \lambda \sigma(t) (1 - \sigma(t)) J(\sigma(t)). \quad [5]$$

The Youden index determines the evolution of the share of Better tenured scientists because it governs the gap between the tenure probabilities of Better and Worse scientists, which determines the difference between the growth rates of the populations of Better and Worse tenured scientists.

The dynamics of the population of scientists follow [5]. Hence, we obtain the following results:

Proposition 1. *Population dynamics depend on power. With high power [$1 - \beta \geq \alpha + \epsilon/(1 + \epsilon)$], the Better paradigm eventually prevails [$\lim_{t \rightarrow \infty} \sigma(t) = 1$], irrespective of initial conditions. With low power [$1 - \beta < \alpha + \epsilon/(1 + \epsilon)$], initial conditions matter: If the initial fraction of Better scientists is high [$\sigma(0) > \sigma^*$], the Better paradigm eventually prevails [$\lim_{t \rightarrow \infty} \sigma(t) = 1$], but if the initial fraction of Better scientists is low [$\sigma(0) < \sigma^*$], the Worse paradigm eventually prevails [$\lim_{t \rightarrow \infty} \sigma(t) = 0$].*

The proof is in [SI Appendix A](#) but is illustrated in Fig. 1. For different levels of power, the Youden index has different properties, modifying the properties of [5]—the differential equation governing the dynamics of the population of scientists.

The proposition provides perspective on the conjecture by Friedman that inferior paradigms will necessarily be abandoned, even if scientific tests have low power (ref. 1, p. 11). If scientists are unbiased ($\epsilon = 0$), since the Better paradigm yields better

of “an excess of [productive capacity] ...beyond the amount required to meet all demands that are backed by the ability and the willingness to pay for the things demanded” (ref. 20, p. 362). Harvard professor Silas Macvane followed the article with a comment that concluded: “Demand for savings is the offer of labor for wages. In order that the supply of capital shall exceed the demand for it, there must be more capital offering for labor than the laborers are willing to receive! The mere statement of the case is sufficient to show its absurdity” (ref. 20, p. 366).

Undaunted, Crocker wrote a book, *The Cause of Hard Times*, in which he developed his underconsumption theory (21). But rather than becoming known as precursor to Keynes, Crocker was ignored. His distress is expressed in the last chapter: “In closing, it may be well to say that no professional economist has ever publicly recognized the validity of the theories and arguments set forth in this book” (ref. 21, p. 103). Among the economists who “published attempted refutations” or who “privately expressed to this author their complete dissent from his views” were luminaries of the profession, including J. Laurence Laughlin, Thorstein Veblen, and Frank Taussig. The Great Depression generated a powerful test of the old paradigm that supply creates its own demand, and, after Keynes’ *General Theory* (22), economists no longer dismissed underconsumption theory as absurd.

The Persistence of Bloodletting. We have seen that power could be low because there are no high-power scientific tests to discriminate between true and false paradigms. But even if such tests do exist, they may play little role in promotion into the fellowship of senior scientists. The history of medicine illustrates this second possibility: High-power tests were available but played no significant role in promotion to the rank of practicing physician. Consequently, some procedures persisted decades after they had been shown to be harmful.

Bloodletting is a good example of such persistence. In the 1830s, Pierre-Charles-Alexandre Louis, a practicing physician in Paris, took matched samples of pneumonia patients: one sample, with bloodletting in the first 4 days of the disease; the other sample, with bloodletting in days five to nine. Louis’ results should at least have called for further testing, since he found a 76% higher fatality rate for those with early treatment (23). That difference was difficult to explain if bloodletting was as beneficial to pneumonia patients as it was thought to be. When published in English (24), Louis’ findings were hailed in the *Journal of the American Medical Society* as “one of the most important medical works of the present century,” being “the first formal exposition of the results of the only true method of investigation in regard to the therapeutic value of remedial agents” (ref. 25, p. 102).

Yet Louis’ use of statistical trials to determine the effects of bloodletting did not catch on. Neither did the later, more conclusive findings by Bennett (26, 27) have significant effect on practice. Bennett found no deaths among 105 patients whom he had treated for pneumonia without bloodletting at the Edinburgh Royal Infirmary; in contrast, when bloodletting had been standard treatment, more than one third of the pneumonia patients had died. The 1909 edition of Osler’s influential textbook on *The Principles and Practice of Medicine* said that “local bloodletting by cupping or leeches is certainly advantageous in robust subjects” (ref. 28, p. 782); such statements remained in posthumous editions, as late as 1942 (29).

Medical historian John Harley Warner culled doctors’ letters and reports to explain why physicians were so averse to statistical methods. He concluded that physicians viewed themselves as professionals with clinical duties toward their patients and that treatment would depend upon physicians’ ability at observation, which was learned through their experience in practice. With this identity, it was considered denial of clinical

duty to base judgment in individual cases on statistical samples of unknown patients in different locales and in different circumstances: “[Doctors] were not prepared to accept even in principle the proposition that they should discard existing therapeutic beliefs and practices, validated by both tradition and their own experience on account of somebody else’s numbers” (ref. 30, p. 201).

In sum, physicians and scientists had different norms regarding promotion. Kuhn’s discussion of the Lavoisier revolution in chemistry illustrates the existence of a scientific norm: Being a scientist entails accepting hypotheses that are confirmed by high-power tests. Lavoisier discovered a superior paradigm about combustion, which replaced an old paradigm that viewed combustion as occurring when flammable materials released their phlogiston. The difference between the two paradigms could be tested through the use of vacuums and precise weights; such tests had high power. Kuhn calls out Lavoisier’s rival, Priestley, for being “unreasonable” because despite findings backed by high-power tests, Priestley resolutely continued his belief in the old phlogiston paradigm; Kuhn even engages in a rare demotion: He judges that Priestley “ceased to be a scientist” (ref. 2, p. 159). In medicine, in contrast, there was no norm that a candidate’s contribution to science should be evaluated with an eye on the results of high-power scientific tests; instead, physicians’ criteria for promotion rested on a candidate’s ability to carry out existing medical practice.

The Persistence of Radical Mastectomy. The history of radical mastectomy further illustrates the resistance of doctors to testing current procedures. Radical mastectomy was introduced in the United States in 1892 by Johns Hopkins’ William Halsted. While the procedure was highly debilitating to its survivors, statistical evidence showed that it was not very effective. Indeed, matched statistics from the Cleveland Clinic published by Crile (31) indicated that radical mastectomy yielded no improvement in mortality relative to simple mastectomy or lumpectomy, which were much less invasive. Despite the evidence, in 1968, 86% of surgical treatments for breast cancer were still by radical mastectomy (ref. 32, p. 132).

After the publication of Crile’s findings, it took more than 10 years before a significant-size randomized controlled trial was begun, against fierce opposition from the cancer-surgeon establishment (32). The opposition continued even when the trial was in progress. The breast-cancer surgeons, like Warner’s 19th-century physicians, based their resistance on their belief in the powers of clinical expertise. In an extreme expression of that opposition, the editor of the journal of the American Cancer Society said that use of randomized controlled trials to decide on procedures for individual patients was playing “scientific Russian roulette” with their lives (ref. 32, p. 115). When its results were published in 1981, 20 years after Crile’s article, the randomized controlled trial bore out Crile’s initial findings: no difference in mortality, but great difference in the condition of the survivors (33).

In medicine, results from high-power tests played no role in promotion to the status of practicing physician. Instead, promotion depended on ability to execute current technique—especially in surgery. For example, trainees in breast-cancer surgery were admitted as practicing surgeons themselves based on their ability to carry out radical mastectomies, so that the promotion to elder of the profession had no regard for Crile’s findings of little difference in mortality but much difference in patient welfare between radical and simple mastectomy.

Conclusion

This paper proposes a model regarding the adoption or non-adoption of superior scientific paradigms. The model captures

