

Drift or Selection: Null Models for Markets

Brecht Corbeel

Absolute Digital Publishers

absolutedigitalpublishers.com

August 30, 2026

Abstract

Population genetics learned in 1968 that a pattern is not evidence of selection until the neutral prediction has been computed, and built exact tests to make that computation routine. Economics has an analogous neutral model, Gibrat's law of proportionate growth, but has never turned it into a test of selection against drift. This paper does so. We define a neutral share process, every firm's expected growth identical and all dynamics noise, and show that under it concentration rises in expectation at a rate fixed by growth-shock variance alone; that the Price selection term and a reallocation-work statistic have exact null moments; and that a lock-in threshold governs when early noise becomes a permanent standard. We build an exact, distribution-free permutation test of a firm's productivity against its subsequent growth, and report its power by simulation: against a ten-percent selection intensity, with growth noise at the Laplace scale of real firm panels, an industry of two hundred firms detects it only two times in five with equal shares and barely more than one time in four under realistic concentration. Concentrated markets have a higher drift barrier: a productivity difference small against $\sigma_\varepsilon\sqrt{\text{HHI}}$ stays invisible to selection no matter how long the market runs.

Keywords: neutral theory; Gibrat's law; genetic drift; permutation tests; market selection; concentration; power laws; effective population size.

JEL codes: B52, C12, C15, D22, L11, L25.

Series: The Evolutionary Economics Papers, No. 10.

1 Introduction

A market-share panel does not come with a p -value attached. A regulator watching an industry's Herfindahl index climb over a decade, an economist finding that a productivity measure predicts subsequent growth with a small but positive coefficient, a case study reporting that the firms which survived a shakeout were, on average, more efficient than the firms which did not — none of these observations is, by itself, evidence that anything was selected. Rising concentration, a positive productivity–growth covariance, and a surviving population fitter on average than the population it replaced are also what a market with no selection at all produces, some of the time, by chance, at a rate that depends on how many firms there are and how noisy their growth is. The question this paper asks is not whether markets select — of

course some of them do, sometimes — but how a reader is supposed to tell a selected pattern from an unselected one when both leave the same kind of footprint in the data. Answering that question requires, before anything else, knowing exactly what the unselected footprint looks like. This paper builds that footprint for markets and turns it into a test.

Population genetics asked the identical question of DNA sequence data in 1968 and spent the following half-century building the machinery to answer it. Motoo Kimura’s founding observation was arithmetic before it was theory: the rate of nucleotide substitution implied by comparing sequences across species “seems to give a value so high that many of the mutations involved must be neutral ones” [68] — selection, on the numbers alone, could not plausibly be driving all of it. King and Jukes reached the same conclusion the following year by a different route, arguing that most amino-acid substitutions accumulate through a process — “non-Darwinian evolution” — that changes nothing fitness can register [70], and Ohta’s nearly neutral extension sharpened the claim into a population-size-dependent one: slightly deleterious mutations fix by drift whenever the population is too small for selection to reliably oppose them [92, 93]. What made the neutral theory a discipline rather than a hypothesis was not that population geneticists agreed with it — many did not, and still do not — but that they built exact statistics whose distribution under strict neutrality could be computed and checked against data: Ewens’s sampling formula for allele partitions [41], Watterson’s independent estimator of the same drift parameter [110], Tajima’s comparison of the two [108], the Hudson–Kreitman–Aguadé test across loci [61], and the McDonald–Kreitman test within one [84]. Lynch’s drift barrier states the resulting limit precisely: selection cannot reliably act on an effect smaller than about the reciprocal of the effective population size, so that whether a given difference is visible to selection at all depends on a population parameter that has nothing to do with the size of the difference itself [80]. Kimura’s lesson, transcribed for any population subject to variation and differential success, is this: a pattern that looks like selection is not evidence of selection until someone has computed what the same pattern would look like without it.

Economics has never lacked a candidate null of its own. Robert Gibrat’s law of proportionate effect — growth uncorrelated with size — is four decades older than Kimura’s paper and was built for exactly the same purpose, to show that a striking regularity (there, the skewed size distribution of firms, cities and incomes) requires no differential fitness to produce [49]. Formalized by Kalecki, Champernowne, Simon and Ijiri [14, 62, 64, 101], surveyed by Sutton [106], and tested against real firm panels by Mansfield, Evans, Hall, Dunne, Roberts and Samuelson, and Lotti, Santarelli and Vivarelli [35, 40, 54, 77, 83], Gibrat’s law and its descendants reproduce the firm-size Zipf law and the Laplace-shaped growth-rate distribution — facts established by Stanley et al., Bottazzi and Secchi, Fu et al. and Axtell [5, 9, 46, 103] — without any firm in the model being better than any other. What economics has never done with this null is what population genetics did with its own: turn it into a statistical test that a reported selection effect must survive. Ecology did the analogous work for community abundance, building Hubbell’s neutral theory [59] into a fitted sampling distribution [109] and then, in the discipline’s model episode, testing it against tropical-forest census data and watching it lose to the ordinary lognormal [85], adjudicated without abandoning the null’s usefulness by Alonso, Etienne and McKane [1]. Section 7 returns to that episode as the standard this paper’s economic test is held to. Economics has the null; it has never built the test, or asked what the test can and cannot see.

This paper builds it. What follows makes five claims, and states now, so that none of it need be inferred later, what each one does and does not assert. (a) The neutral share

process — firms whose shares evolve by $s_i(t+1) \propto s_i(t)(1 + \varepsilon_i)$ with ε_i iid and mean zero — is stated as the economic analogue of the Wright–Fisher process, and Proposition 1 derives its consequence for concentration: under pure noise, the Herfindahl index rises in expectation at a rate fixed by the growth-shock variance, so a concentrating industry is not, on its own, an industry that is selecting anyone. (b) That same null is shown to be the one already scattered, largely without being named as one, through three companion papers in this series: the drift benchmark the Price selection term needs [30], the drift benchmark for reallocation work under creative destruction [29], and the lock-in threshold of a positive-feedback standards race [23]. Proposition 2 assembles the three into one framework carrying a single parameter pair, the growth-shock variance and the effective population size $N_e = 1/\text{HHI}$. (c) Proposition 3 turns the assembled null into an exact test: under the null hypothesis that growth is independent of a firm’s productivity, the joint distribution of growth rates is invariant to permuting the productivity labels across firms, so the exact null distribution of the Price selection term — and of any statistic built the same way — is obtained by recomputing it under repeated relabelings, exactly as Fisher’s original randomization logic prescribes [42]. (d) Proposition 4 reports the test’s power by simulation and states the drift barrier this paper’s title promises: a productivity difference is invisible to market selection once its effect on growth is small against the drift scale set by growth noise and concentration, transcribing Ohta’s and Lynch’s population-genetic drift barrier into a market’s diversification and turnover statistics [80, 92]. (e) Section 7 re-reads a set of already-published, already-verified findings on weak market selection through this null, asking at each one not whether the reported effect is real but whether the study’s own sample size and industry concentration could, at conventional power, have told a real effect of that size apart from drift.

None of this is a claim that any particular industry is, or is not, selecting on efficiency; it is a claim about what a reader needs to compute before believing either answer. §2 sets out the population-genetics discipline this paper transcribes, table by table. §3 does the same for the economic neutral model, tabulating its formalizations, its tests and the distributional facts it explains. §4 states the neutral share process and Propositions 1 and 2 in full. §5 states Proposition 3, the five-step permutation protocol, and what is and is not testable this way. §6 reports Proposition 4’s power surface and the drift barrier. §7 re-reads the empirical selection literature against it, following the ecological precedent of §2. §9 states what the null changes for the papers that came before it in this series, proposes the reporting standard the permutation test makes possible, and closes.

2 The neutral theory and the discipline it created

In 1968, Motoo Kimura proposed that calculating the rate of molecular evolution in nucleotide substitutions “seems to give a value so high that many of the mutations involved must be neutral ones” [68]. Working independently, King and Jukes argued the same case the following year under the banner of “non-Darwinian evolution”: most of the amino-acid and nucleotide substitutions accumulating in evolving lineages are, on this reading, invisible to natural selection because they change nothing that fitness can register [70]. Kimura spent the following fifteen years building the claim into a full theory, resting on the diffusion-equation treatment of fixation probability he had already worked out in 1962 [67] and codified in his 1983 book [69]. Tomoko Ohta’s nearly neutral extension sharpened the claim in the direction this paper needs most: mutations that are only slightly deleterious, not perfectly neutral, are still fixed by

drift once the population is small enough that selection cannot reliably oppose them, so that “evolution is rapid in small populations or at the time of speciation” [92], a population-size dependence Ohta later generalized into the nearly neutral theory proper [93]. None of these papers denies that selection exists. What they assert is that a great deal of what looks, on inspection, like the signature of selection is the signature of population size and chance, and that telling the two apart is a job for a statistical test, not for a plausibility argument.

Building that test took population genetics several further steps, each supplying an independent handle on the same null. Ewens’s sampling formula gives the exact probability, under the neutral infinite-alleles model, of any observed partition of a sample of genes into allelic classes, as a function of sample size and a single mutation-scaled parameter [41]. Watterson supplied a second, independent route to that same parameter, estimated from nothing but the number of segregating sites in a sample of DNA sequences [110]. Because the identical parameter can also be recovered from the average number of pairwise differences between sequences — a frequency-weighted rather than a site-counting quantity — strict neutrality makes a falsifiable prediction that the two estimators agree, and Tajima built his test directly on that prediction: a statistic, since bearing his name, formed from “the relationship between the two estimates of genetic variation at the DNA level, namely the number of segregating sites and the average number of nucleotide differences,” constructed so that it is centered at zero under strict neutrality and constant population size and displaced from zero by population growth, contraction, structure, or selection [108]. The verification behind this paper could not independently confirm Tajima’s own defining formula from the primary 1989 text — the source PDF resisted every access route attempted — so the description kept here is the qualitative one on which the field agrees, not a quotation from the source. Hudson, Kreitman and Aguadé generalized the comparison from one locus to several: under neutrality the ratio of within-species polymorphism to between-species divergence should be constant, up to a locus-specific mutation rate, across every locus sampled, and the resulting joint statistic across loci — the HKA test — found, applied to the *Adh* region of *Drosophila*, a departure from that constancy [61]. McDonald and Kreitman closed the sequence with a test needing only a single locus and no assumption about population size at all: classify substitutions as replacement or synonymous, classify each further as a polymorphism within species or a fixed difference between species, and neutrality requires the same replacement-to-synonymous ratio in both classifications. At the same *Adh* locus, “there are more fixed replacement differences between species than expected” [84], an excess the authors read as adaptive protein evolution rather than drift. Four tests, one method: state exactly what a sample should look like if nothing but population size and mutation were operating, and see whether the data comply.

Building the tests did not end the argument; it relocated the argument onto the tests’ own results, and the relocation is instructive because economics has no equivalent argument to point to. Nei’s survey of the selectionism/neutralism controversy shows a field that never fully closed the question, proposing that redefining neutrality functionally — mutations that do not appreciably change what a gene product does — absorbs much of the dispute without eliminating a role for adaptive change [89]. Kreitman’s and Nielsen’s review-length treatments catalog the resulting family of neutrality tests and the assumptions each leans on, with Kreitman stating the shared logic plainly: “all tests use predictions from the theory of neutrally evolving sites as a null hypothesis” [72, 91]. Hahn’s 2008 paper pressed the selectionist side explicitly, tracing the mathematical lineage of neutral null models from Fisher and Wright through Kimura and arguing for giving selection a more central explanatory role than strict neutral nulls allow

[53]. Kern and Hahn’s fiftieth-anniversary reassessment went further still, concluding that the neutral theory “lacks empirical support and should be rejected” in favor of genomes “shaped in prominent ways by the direct and indirect consequences of natural selection” [65]. Seven co-authors, including Lynch and both Charlesworths, replied in the same journal the following year to “critically examine and ultimately reject Kern and Hahn’s arguments and assessment,” concluding instead that “it is now abundantly clear that the foundational ideas presented five decades ago by Kimura and Ohta are indeed correct” [63]. This paper takes no side in that exchange and needs to take none: what matters for the argument here is not who is right about molecular evolution but that the discipline conducts the argument at all, in public, using the tests of the previous paragraph as the shared instrument, rather than trading unquantified intuitions about how much of a pattern “feels like” selection. Gillespie’s graduate text is the standard record of the debate conducted on those terms, weighing selectionist and neutralist readings against the same body of test statistics rather than against each other’s priors [50].

The same body of theory supplies a second result this paper borrows directly: a lower bound on what selection can see at all. Lynch states it as a criterion on the fixation probability of a mutation with selection coefficient s in a population of effective size N_g : fixation is governed by $S = 2N_g s$, “twice the ratio of the power of selection to the power of random genetic drift,” so that “if $1/N_g \gg |s|$, the population will be driven to a state expected under mutation pressure alone” [80]. Selection has a floor set by population size, not by how large or small s is in some absolute sense: an effect decisive in a small population is invisible in a large one. Lynch’s companion book develops the same criterion at length as the drift barrier governing genome architecture [81], and Lynch and Conery’s empirical companion argues that much of the increase in genome complexity across the tree of life is better read as a passive consequence of long-run reductions in effective population size than as adaptation [79]. Charlesworth’s review states plainly what stands behind N_g in every one of these results: effective population size N_e is “a core concept in population genetics ... necessary for determining the evolutionary role of genetic drift,” fixing simultaneously the level of neutral variation a population carries and “the effectiveness of selection relative to drift” [15]. Section 6’s drift barrier for markets, and the effective population size $N_e = 1/\text{HHI}$ that recurs through every proposition of this paper, are the direct economic transcription of this pair of results; Appendix B states the transcription exactly.

A second, independent lineage carried the same logic from molecules to whole communities, and its history matters here because it goes one step further than population genetics ever needed to: it produces a documented case of the null being fit against real field data and losing. Hubbell’s 1979 study of a tropical dry forest already reads relative species abundance through “a simple stochastic model based on random-walk immigration and extinction set in motion by periodic community disturbance” [58] rather than through niche differences among species. His 2001 book, *The Unified Neutral Theory of Biodiversity and Biogeography*, generalizes the same drift logic — every individual of every species subject to identical per-capita birth, death, dispersal and speciation rates, regardless of species identity — into a general theory of community assembly [59]. Bell’s independent and contemporaneous statement of “neutral macroecology” shows the same family of abundance and diversity patterns able to “arise...from neutral community models in which all individuals have identical properties, as the consequence of local dispersal alone,” concluding that “functional interpretations of these patterns must be reevaluated” [8]. Chave’s review situates both against the population-genetics lineage explicitly: the defining assumption, equivalence among individuals, “finds little empirical support in general,”

though “a weaker assumption holds for stable communities, the equivalence of average fitness among species” [16], and Leigh’s historical account draws the line directly from Kimura’s “pans-electionist paradox” to Hubbell devising “a similar neutral theory for forest ecology,” arguing that in both settings the null’s value lies in its capacity to “provide probing null hypotheses” rather than to describe nature exactly [74]. Volkov, Banavar, Hubbell and Maritan gave the theory its own testable sampling distribution, deriving the zero-sum multinomial abundance distribution analytically and fitting it to tropical-forest census data [109] — the ecological counterpart, in effect, of Ewens’s sampling formula.

That testable distribution is what let the null be beaten. McGill fit the zero-sum multinomial against the ordinary lognormal on the very data Hubbell’s theory was built to explain — North American Breeding Bird Survey replicates and the Barro Colorado Island tropical tree census — and reported the result without qualification: “Here I test whether the ZSM fits several empirical data sets better than the lognormal distribution. It does not.” For Barro Colorado specifically, across every one of eight goodness-of-fit measures computed, “the lognormal beats the ZSM on all measures of goodness-of-fit,” and more generally, “the ZSM fail[s] to fit empirical data better than the lognormal distribution 95% of the time” [85]. This is the discipline’s model result for the present paper, returned to directly in the re-reading of Section 7: a null taken seriously enough to be fit, quantified, and beaten on its own terms by real data, rather than argued against on priors. Alonso, Etienne and McKane’s response supplies the qualification that keeps the episode a model of discipline rather than an overreaction: a failed goodness-of-fit test is not, by itself, a rejection of the theory’s mechanism, because “neutral theory is also a stochastic theory, a sampling theory and a dispersal-limited theory,” and each component can be probed and revised independently of the others [1]. Etienne’s improved multivariate sampling formula is a product of exactly that kind of revision, sharpening the statistical machinery used to fit the neutral abundance distribution without abandoning the neutral framework [39], and Hubbell’s own later reflection on the “hypothesis of functional equivalence” — the claim that trophically similar species are “demographically identical on a per capita basis” — is explicit that this is the falsifiable core the whole edifice stands or falls on [60]. Rosindell, Hubbell and Etienne’s ten-years-on retrospective states the field’s resulting self-understanding directly: neutral theory’s contribution is that it “provides quantitative null models for assessing the role of adaptation and natural selection,” not a claim that communities are actually neutral [98].

Both lineages are instances of a broader methodological genre that ecology named and codified explicitly. Gotelli and Graves’s reference text collects the randomization and Monte Carlo procedures — reshuffle a co-occurrence matrix, redraw a community from a regional pool, hold row and column totals fixed and ask what pattern survives — that operationalize the comparison of observed data against a specified stochastic process across community ecology and biogeography [52], and Harvey, Colwell, Silvertown and May’s earlier survey establishes the same discipline as a general research program rather than a set of isolated tricks: before reading any pattern as evidence of a biological process such as competition, construct the null model of what pattern would appear in its absence [56]. Read together with the population-genetics tests above and the ecological drift models and their adjudication, the two literatures amount to a single working method, applied twice over some sixty years to two different objects: state the null exactly enough to compute what it predicts, build a statistic whose distribution under that null is known, and be willing, as Barro Colorado Island shows, to lose. Table 1 states, biological test by biological test, what this paper’s assembled null (§4) makes the corresponding

Biological test	What neutrality predicts	Economic analogue
Ewens sampling formula / Watterson's θ	the exact distribution of allele partitions, and an independent estimator of the same drift parameter from segregating sites	the distribution of HHI under the neutral share process (Proposition 1), and $N_e = 1/\text{HHI}$ as the market's diversity estimator
Tajima's D	a standardized difference between two diversity estimators, centered at zero under strict neutrality	$T_{\mathcal{W}}$: the standardized difference between realized reallocation work $\mathcal{W}$ and its neutral expectation $\frac{\sigma_\varepsilon^2}{2}(1 - \text{HHI})$
Hudson–Kreitman–Aguadé (HKA)	joint constancy of the polymorphism/divergence ratio across several loci	joint constancy of the selection term across several industries, combined by Fisher's or Stouffer's method
McDonald–Kreitman	a 2×2 comparison of replacement/synonymous substitutions against polymorphism/divergence, with an excess of replacement fixation reading as selection	the permutation test of Sel on the pairing of z with (s, w) , with an excess of covariance beyond the permutation null reading as selection
The drift barrier (Ohta; Lynch)	selection cannot see an s smaller than about $1/N_e$	market selection cannot see a δz smaller than about $\sigma_\varepsilon \sqrt{\text{HHI}}/\beta$; concentrated markets have a higher floor

Table 1: Biological neutrality tests and their economic analogues, as assembled from the null of §4. Appendix B states each analogue's definition in full.

economic object; Appendix B gives each analogue's exact definition.

3 The economic neutral model

Population genetics did not invent the idea that a system's most visible statistical regularities can arise without anyone being selected. Economics had its own version four decades before Kimura, in a single monograph on the statistics of firms, cities and incomes: Robert Gibrat's *loi de l'effet proportionnel*, the law of proportionate effect [49]. Gibrat's claim is that a firm's proportional growth rate is drawn independently of its current size — large and small firms grow, in mean and in dispersion, at the same rate. Iterated over many periods this generates a lognormal cross-section: log size is a sum of many small, independent, size-blind shocks, and the central limit theorem does the rest. No firm in Gibrat's model is better than any other. The size distribution it produces — skewed, long right tail, occasional giants — looks like a century of industry data, for a reason that has nothing to do with efficiency, management quality, or survival of the fitter. This is economics' neutral model in Kimura's exact sense: a generative account of an observed pattern under the hypothesis that the variation across units is functionally irrelevant to their fate.

Gibrat's law took thirty years to grow from one claim into a family of models sharing that neutral spine. Kalecki showed, in the paper that gave his name to the “Gibrat distribution,” that the naive reading is self-undermining: if the variance of log size grows linearly with firm age, as strict iid proportionate shocks imply, a cross-section of firms of different ages cannot converge to one stationary lognormal without a stabilising ingredient — a floor on size, or a growth-size

correlation the strict law rules out [64]. Champernowne’s fully specified Markov process on proportional income classes produced not a lognormal body but a Pareto upper tail, moving the family from “proportionate growth explains the shape” to “proportionate growth explains *which* shape, tail included” [14]. Simon derived a closely related tail by an entirely different route — new individuals (words, papers, cities, incomes, firms) assigned to existing classes in proportion to each class’s current size — as the stationary form of the Yule distribution, noting that the identical mechanism reproduces word frequencies, city populations, incomes and biological genera all at once [101]: proportionate, share-weighted addition is a mechanism-agnostic route to skew distributions, not a fact peculiar to firms. Simon and Bonini fitted the machinery directly to firm growth, adding an entry process to the recursion and showing the combination reproduces the observed right-skewed, approximately lognormal firm-size distribution [102], and Ijiri and Simon collected the resulting research programme into one account of skew firm-size distributions [62]. Two further threads closed the loop between theory and data: Hart and Prais ran one of the first large-sample statistical tests of the lognormal, proportionate-growth model against UK business-concentration data [55], and Steindl built an explicit birth-growth-exit process for firms whose steady state is a Pareto law — the first place the family produces a power-law tail from entry and proportionate growth rather than a Markov chain on income classes [104]. By the mid-1960s Gibrat’s law was a cluster of stochastic-process models sharing one neutral core — growth uncorrelated with the characteristic modelled — differing only in which boundary condition turns that core into a lognormal body, a Pareto tail, or both.

Testing whether real firms obey this core claim took another round of work, and the verdict, tabulated by Santarelli, Klomp and Thurik, is that mostly they do not — not at the small end, and not in the variance. Mansfield’s tests on U.S. steel, petroleum-refining and rubber-tire manufacturers across several ten-year windows between 1916 and 1957 rejected strict Gibrat’s law in most of ten formulations: seven of ten in the static comparison of growth-rate distributions across size classes, with smaller firms markedly more likely to exit, and a further four of ten and three of ten in a distributional and a regression variant, growth-rate variance itself falling with size in a fourth [83]. Evans, using some twenty thousand U.S. manufacturing firms over 1976–1982, found growth declining with size and the departure from Gibrat’s law concentrated among the smallest and youngest firms, consistent with Jovanovic’s learning account — the magnitude is not recoverable here, only the direction and its concentration at small size [40]. Hall, on 1,778 large, publicly traded U.S. manufacturers, 1972–1983, found the same negative size–growth relation holding almost uniformly across the size range, with growth-rate variance again declining with size [54]. Dunne, Roberts and Samuelson, tracking 219,754 U.S. manufacturing plants entering in 1967, 1972 or 1977 through 326,936 plant-year records, rejected Gibrat’s law for surviving and all plants alike, with failure risk and subsequent growth both falling in size and age [35]. Lotti, Santarelli and Vivarelli’s study of newborn Italian firms supplies the reconciling pattern, restated by the same authors in 2007: Gibrat’s law “has to be rejected *ex ante*, since smaller firms tend to grow faster than their larger counterparts,” but the same cohorts show “a significant convergence towards Gibrat-like behavior” *ex post* [77]. Sutton’s survey and Coad’s later synthesis reach the same qualified verdict: strict, size-independent, constant-variance Gibrat growth is not a literal description of firm dynamics at entry, but an approximate regularity the surviving population relaxes toward once selection at the smallest sizes has run its course [19, 106]. None of this counts against using Gibrat’s law as a null model for markets generally; it says where the null is expected to fail — at entry, where survival itself is at stake — and where it holds well once that margin is conditioned away. Two

later panels sharpen this age-conditioned picture rather than overturning it. Coad, Segarra and Teruel’s study of Spanish manufacturing firms finds ageing associated with steadily rising productivity, profits, size and equity ratios but with lower expected growth rates of sales, profits and productivity — performance improves with age on some measures and worsens on others, not uniformly in either direction [20]. Dosi, Grazzi, Moschella, Pisano and Tamagni’s tracking of US manufacturers from 1959 to 2015 finds that sustained high growth, where it occurs, “is there and may be even substantial for some firms, but it is rare” [34]: the long-run growth persistence a strong-selection account would predict is the exception rather than the norm, consistent with a population converging toward Gibrat-like behavior once the entry margin is conditioned away.

The size-and-growth facts get sharper once physicists start measuring them alongside economists. Stanley et al., using Compustat data on the whole population of publicly traded U.S. manufacturers from 1975 to 1991, found the one-year log growth-rate distribution for firms of a given size to be tent-shaped on a log scale — exponential, not Gaussian — with its standard deviation falling as a power law in firm size over roughly seven orders of magnitude [103]. Amaral et al. pinned the exponent down on the same Compustat universe: the width of the growth-rate distribution scales as $S^{-\beta}$ with $\beta = 0.20 \pm 0.03$ for sales and for property, plant and equipment, and $\beta = 0.18 \pm 0.03$ for employees, assets and cost of goods sold [2]. Bottazzi and Secchi confirmed the tent shape is no Compustat artifact: fitting the flexible Subbotin family to pooled growth rates across fifty-five Italian manufacturing sectors, the maximum-likelihood shape parameter is $b = 0.965$ (s.e. 0.007), indistinguishable from the Laplace value $b = 1$ in forty of the fifty-five sectors, and a neutral “islands” model of opportunity assignment — Bose–Einstein rather than the equiprobable, binomial assignment of the classical Simon-type island models, which instead yields a Gaussian growth density — turns a population of size-blind, proportionate individual shocks into exactly this heavy-tailed, tent-shaped aggregate [9]. Fu et al. give the growth-rate density its full shape across several levels of aggregation — products, firms, countries — a Laplace cusp at the centre and a power-law tail with exponent $\zeta = 3$ [46]. On the cross-sectional size distribution, Axtell’s census of the entire U.S. firm population found a Zipf law with exponent close to one: $\alpha = 1.059$ from a full-population regression, $\alpha = 1.098$ (s.e. 0.064) restricted to firms with employees using 1992 Census data, and a run of year-by-year estimates over a ten-year period clustering between 0.9986 and 1.0039 despite the number of firms, total employment and average firm size all changing over the decade [5]. Cabral and Mata, tracking a comprehensive panel of Portuguese manufacturing firms, showed cohorts are born right-skewed and evolve toward lognormal as they age — but that “selection accounts for very little of this evolution,” with financing constraints doing most of the explanatory work [12]. The tent shape, its size-dependent width, the Laplace centre with a power-law tail, the near-exact Zipf cross-section, the cohort’s drift toward lognormality — each is exactly what proportionate, size-blind growth with an entry floor produces, with no firm fitter than any other.

That a neutral process can manufacture a Zipf law, and not merely a lognormal body, is a mechanical fact with its own paper trail predating Gibrat by six years. Yule showed that a branching process in which new species are added to genera in proportion to each genus’s existing size generates the same right-skewed, heavy-tailed shape seen in species-per-genus counts [111] — the ancestor of every preferential-attachment account of skew distributions since. Kesten’s theorem supplies the general mechanism: a random multiplicative recursion with an additive term, $X_{t+1} = A_t X_t + B_t$, converges to a stationary distribution with a power-

law tail even when neither A_t nor B_t is itself power-law distributed [66]. Gabaix specialised this to cities: ordinary Gibrat growth combined with nothing more than a lower reflecting bound converges to an exact Zipf law with exponent one — the canonical demonstration that pure drift plus a floor, with no optimisation or selection, is sufficient [47], a result he later placed in a wider survey of power-law-generating mechanisms across economics and finance [48]. Newman’s tutorial catalogues the equivalent routes — multiplicative processes, preferential attachment, critical phenomena, random walks with a boundary — converging on the same functional form [90], and Reed shows one more: geometric Brownian motion, the continuous-time analogue of Gibrat growth, observed at a random, exponentially distributed stopping time produces a Pareto tail on its own, no floor required [97]. Levy and Solomon make the underlying point almost trivial: take logarithms of a multiplicative process and it becomes an ordinary additive Boltzmann process, so a power law in levels is nothing more than an exponential law in log space [75], and Sornette and Cont supply the companion convergence result for a multiplicative process repelled from zero [21]. Malevergne, Saichev and Sornette reframe Zipf’s law as the condition under which incumbents’ size-weighted growth balances the growth contributed by new entrants — a sustainable-growth constraint rather than an arbitrary drift outcome [82], and Saichev, Malevergne and Sornette’s monograph ties the whole Kesten-type family into one synthesis [100]. Arata’s study of over one hundred thousand Japanese firms supplies the needed qualification: a Laplace-shaped growth distribution requires occasional large discrete jumps, not merely many small multiplicative shocks [4] — continuous drift alone is not quite enough, though drift-with-jumps still needs no fitness differential between firms.

Mitzenmacher’s history of these generative models states the problem this paper is built to solve: power-law and lognormal generators are close mathematical cousins — multiplicative processes with or without a reflecting floor or an optional-stopping rule — and are easily confused from finite data [88]. Two case studies show how easily. Luttmer builds a model of idiosyncratic productivity shocks, selection of the successful, and imitation by entrants that reproduces Zipf’s law for firm sizes just as well as any pure-drift account, calibrated to attribute roughly half of aggregate output growth to that selection [78]; Rossi-Hansberg and Wright build an investment-in-human-capital account of establishment size dynamics matching the same aggregate growth, exit and size-distribution facts through an entirely different, selection-flavoured mechanism [99]. A Zipf exponent near one is consistent with pure neutral drift and with two structurally different selection stories; the size distribution alone cannot adjudicate between them. Eeckhout and Levy’s exchange over city sizes illustrates the same trap at the level of the tail itself. Eeckhout, using the full, non-truncated 2000 Census city-size distribution, found it lognormal rather than Pareto, reconciled with ordinary proportionate growth [36]; Levy’s comment showed the top 0.6% of cities, over 23% of the urban population, trace a shape “dramatically different from the lognormal,” and Eeckhout’s reply countered that the visual log-log reading of a truncated tail is “misleading” and the lognormal survives “tail included” [76]. Two careful economists, looking at the same data, disagreed about which distributional family it belongs to — exactly the point at which eyeballing a log-log plot stops being an adequate test.

That disagreement is what Clauset, Shalizi and Newman built a statistical answer to. Maximum-likelihood fitting, a Kolmogorov–Smirnov goodness-of-fit test and a likelihood-ratio comparison against alternatives, applied to twenty-four datasets each previously claimed to follow a power law, found seventeen consistent with the power-law model and firmly ruled it out for the remaining seven [17]. Stumpf and Porter’s companion note turns the same discipline

Table 2: Empirical tests of Gibrat’s law and the size and growth distributions it predicts.

Study	Sample	Finding	Source
Hart–Prais 1956	UK quoted manufacturers, large sample	Early large-sample support for the lognormal, proportionate-growth concentration model	[55]
Mansfield 1962	US steel, petroleum-refining, rubber-tire firms, 1916–1957	Gibrat’s law rejected in most of ten test formulations; rejection concentrated among smaller, exiting firms	[83]
Evans 1987	~20,000 US manufacturing firms, 1976–1982	Growth declines with size, most among small/young firms; consistent with Jovanovic learning	[40]
Hall 1987	1,778 large US public manufacturers, 1972–1983	Negative size–growth relation, fairly uniform across size; growth-rate variance falls with size	[54]
Dunne–Roberts–Samuelson 1989	219,754 US plants entering 1967/72/77	Gibrat’s law rejected for surviving and all plants; failure and growth both decline in size and age	[35]
Lotti–Santarelli–Vivarelli 2003	Newborn Italian firms, whole industrial spectrum	Rejected ex ante (smaller firms grow faster); converges to Gibrat-like behavior ex post	[77]
Cabral–Mata 2003	Portuguese manufacturing panel	Cohorts born right-skewed, evolve toward lognormal; selection explains very little of the evolution	[12]
Axtell 2001	Entire US firm population, Census data	Zipf law, $\alpha \approx 1.0$ –1.1, stable across a decade of changing scale	[5]
Bottazzi–Secchi 2006	55 Italian manufacturing sectors	Growth-rate density Laplace-shaped, Subbotin $b = 0.965$; indistinguishable from $b = 1$ in 40/55 sectors	[9]

into a general warning: most reported power laws in the literature lack either the statistical support or the mechanistic backing to be taken at face value [105], and Perline sharpens the diagnosis by naming the failure mode — a “strong” inverse power law holds over the full data range, a “weak” one only in an upper tail, and a “false” one is a truncated exponential-type distribution that merely looks linear over the limited range a finite sample happens to cover [95]. This is the testing standard the rest of this paper borrows: a distributional claim about firm sizes or growth rates is not established by a log-log plot that looks straight, any more than a selection claim about market shares is established by a Herfindahl index that happens to rise. Table 2 collects the empirical Gibrat-law tests surveyed above; §4 states, in one notation, the neutral share process every one of them was testing against, and §5 extends Clauset–Shalizi–Newman’s distributional discipline from the shape of the size distribution to the selection question itself.

4 One null for markets

The neutral theory of §2 and the Gibrat-law literature of §3 are, formally, the same claim made twice: draw a growth increment independently of everything that might make one unit fitter than another, iterate, and see what the resulting population looks like. What Papers No. 1,

3 and 7 of this series each built, separately, to answer a different question — how much of measured productivity growth is reallocation [30], how much churn a market generates for free [29], how decisively an installed base freezes an accident [23] — turns out to be three questions asked of one process. This section writes that process down once, states what it implies for concentration, and shows that the three papers’ separate results are three readings of a single effective population size.

Definition 1 (The neutral share process). *Firm shares evolve as*

$$s_i(t+1) = \frac{s_i(t) w_i}{\sum_j s_j(t) w_j}, \quad w_i = 1 + \varepsilon_i,$$

with ε_i iid across firms i and periods t , mean 0 and variance σ_ε^2 , independent of every characteristic z_i .

This is Gibrat’s law of proportionate growth [49] written for a closed population — shares that must sum to one, so growth is relative rather than absolute — in the form of a Wright–Fisher-type share process; it is the economic neutral model in exactly Kimura’s sense [68]: every firm has the same expected growth factor, $\mathbb{E}[w_i] = 1$, and everything that follows is noise. Three consequences follow directly, and they are the ones the rest of this paper needs. First, growth rates are size-independent: nothing in w_i ’s law refers to $s_i(t)$, which is Gibrat’s own defining property, stated once and inherited by the whole process. Second, the share vector performs a random walk on the simplex $\{s : s_i \geq 0, \sum_i s_i = 1\}$, driven by the same shocks that drive w , with no restoring force pulling any firm’s share back toward where it started. Third, the size distribution of a cohort that starts together and is tracked forward tends toward the lognormal [12, 64], and once entry and a lower reflecting barrier are added to the pure process — a firm cannot shrink below some minimum viable size without exiting — the stationary cross-section acquires a power-law tail [47, 66, 78]. This is why §3’s Zipf-law evidence [5] is not, by itself, evidence against Definition 1: a reflecting floor on an otherwise driftless random walk produces a Zipf tail on its own, with no firm doing anything a rival could not equally have done.

Proposition 1 (Concentration drifts upward). *Under the neutral share process, to second order in ε ,*

$$\mathbb{E}[\Delta \text{HHI}] = \sigma_\varepsilon^2 \left(\text{HHI} - 4 \sum_i s_i^3 + 3 \text{HHI}^2 \right),$$

which is positive whenever shares are dispersed (for N equal shares it equals $\sigma_\varepsilon^2(N-1)/N^2$) and negative only near monopoly, where the index is bounded above. Appendix A derives the formula from the second-order expansion of $\sum_i s_i^2(1 + \varepsilon_i)^2/(1 + \bar{\varepsilon})^2$.

Corollary 1 (Rising concentration is not evidence of selection). *A neutral industry — one in which every firm has the same expected growth and nothing else is true of it — concentrates in expectation, at a rate set by the growth-rate variance σ_ε^2 and the industry’s own starting dispersion, and the effective population size $N_e = 1/\text{HHI}$ of Paper No. 1 [30] falls under pure noise. A rising Herfindahl index is therefore, on its own, exactly the evidence Proposition 1 predicts drift alone will manufacture, and cannot be read as evidence that better firms are winning until the noise-only path has been subtracted from it.*

The mechanism is Gibrat’s own variance-instability problem [64] read off in a single statistic: a lucky run for one firm and an unlucky one for a similarly sized rival do not cancel in their effect on $\sum_i s_i^2$, they add to it, because squaring is convex. Because N_e is falling exactly when HHI is rising, Proposition 1 says a neutral market’s own effective population size shrinks on average under pure noise, without any firm having done anything — the market analogue of Charlesworth’s point that N_e is itself a quantity driven by the dynamics of drift, and not only a fixed backdrop against which selection is measured [15]. Figure 1 shows the formula holding to close numerical precision: from an equal-shares start ($\text{HHI}_0 = 0.0100$) and from a Zipf-like start ($\text{HHI}_0 = 0.0188$), the simulated one-period $\mathbb{E}[\Delta\text{HHI}]$ matches Proposition 1’s value almost exactly, and over forty periods of nothing but noise the mean Herfindahl index rises from 0.0100 to 0.0148 in the first economy and from 0.0188 to 0.0279 in the second — each acquires half again its starting concentration, and not one firm in either economy is better than any other.

The parameter carrying this result, $N_e = 1/\text{HHI}$, is not new to this paper: it is Paper No. 1’s effective population size [30], imported from population genetics’ use of an effective size to set the strength of drift relative to selection [15, 80]. What Proposition 1 adds is that N_e is not merely a parameter fixed at the start of an episode and read off at the end; it has its own dynamics under the neutral process, falling in expectation whenever shares are dispersed at all. §6 returns to this fact as the drift barrier of concentrated markets. Here it is enough to note that the same N_e reappears, unchanged in its definition, in each of the three results assembled next.

Proposition 2 (The assembled nulls). *Under the neutral share process:*

- (i) *the Price selection term has $\mathbb{E}[\text{Sel}] = 0$ and, to first order, $\text{Var}[\text{Sel}] = \sigma_\varepsilon^2 \sum_i s_i^2 (z_i - \bar{z})^2$ (Paper No. 1 [30]);*
- (ii) *the reallocation work has $\mathbb{E}[\mathcal{W}] \approx \frac{\sigma_\varepsilon^2}{2} (1 - \text{HHI})$ (Paper No. 3 [29]);*
- (iii) *with positive feedback added, the limiting share of a standard is Beta-distributed with variance $p(1-p)/(N_0 + 1)$, and the lock-in number N_{ek} decides whether early noise is decisive (Paper No. 7 [23]).*

Corollary 2 (One null, one dial). *The three results of Proposition 2 are one null with one parameter pair $(\sigma_\varepsilon, \text{HHI})$ plus, where feedback exists, k ; every selection claim made about market-share data is a claim that the data lie outside what this pair generates.*

Read the three side by side against that single pair. Part (i) says the Price covariance between growth and productivity [96] is centered at zero under neutrality — unsurprising, since w and z are independent by construction — but its spread is not arbitrary: it scales with σ_ε^2 and with how unevenly productivity is distributed across share-weighted firms, so a market with a wide productivity spread will show larger neutral selection-term swings than a homogeneous one at the same σ_ε , purely from noise. A selection term reported without its own null spread cannot be told from this. Part (ii) says the churn a market must show every period, from noise alone, has a floor of $\sigma_\varepsilon^2(1 - \text{HHI})/2$: fragmented markets, HHI near zero and N_e large, generate almost the maximum reallocation a given noise level allows, while concentrated markets, HHI near one and N_e near one, generate almost none, because $1 - \text{HHI}$ is exactly the probability that two share-weighted draws land on different firms — the fraction

Statistic	What the neutral share process predicts	Source
Cohort size distribution	Lognormal in the short-to-medium run	[12, 64]
	Power-law tail (Zipf, $\alpha \approx 1$) once entry and a reflecting floor are added	[5, 47]
Growth rates	Size-independent mean; exponential (Laplace) center with a power-law tail in dispersion	[9, 46, 103]
Concentration	$\mathbb{E}[\Delta\text{HHI}] = \sigma_\varepsilon^2(\text{HHI} - 4 \sum_i s_i^3 + 3\text{HHI}^2)$, positive except near monopoly; $N_e = 1/\text{HHI}$ falls in expectation	Proposition 1
Selection term	$\mathbb{E}[\text{Sel}] = 0$; $\text{Var}[\text{Sel}] = \sigma_\varepsilon^2 \sum_i s_i^2 (z_i - \bar{z})^2$	Proposition 2(i); [30]
Reallocation work	$\mathbb{E}[\mathcal{W}] \approx \frac{\sigma_\varepsilon^2}{2}(1 - \text{HHI})$	Proposition 2(ii); [29]

Table 3: What the neutral share process of Definition 1 predicts for the statistics an empirical selection study reports, with no selection anywhere in the generating process.

of the economy that is even eligible to be reallocated between. Part (iii) says that once a positive-feedback mechanism is grafted onto the neutral process — market share begetting more market share, the mechanism by which a technology standard locks in — the outcome stops being a moving distribution and becomes a single realized draw, Beta-distributed with a variance set by the same kind of effective size, N_0 , that Papers No. 1 and 3 call N_e ; the lock-in number $N_e k$, the product of the effective population size and the feedback strength, is the single dial deciding whether that draw is close to deterministic or close to a coin toss. The same variance that drives concentration upward in Proposition 1 sets the spread of the selection term in (i) and the size of the churn floor in (ii), and the same N_e that measures how concentrated a market already is sets, once multiplied by a feedback strength, how much of its future is already decided in (iii). Nothing in any of the three needs a firm to be better than any other.

Reading market selection through a Price covariance has direct precedent in evolutionary economics rather than being an import invented for this paper. Andersen applies Price’s equation directly to economic evolution, partitioning observed change in a population of firms or industries into a selection effect and what he calls an innovation effect [3], and Metcalfe formalizes market competition among firms as a replicator process obeying a Fisher’s-Principle-like relation between the rate of change of a population’s mean characteristic and the variance of that characteristic within it [87] — the bridge from the covariance formalism of Part (i) to competition among firms specifically. Knudsen argues that economic selection theory should be built to mimic the causal structure of its neo-Darwinian counterpart — variation, heredity, and selection as separable stages — rather than borrow the vocabulary of selection without the structure [71], and Frank’s own restatement of the Price equation, addressing critiques of how far its selection/non-selection partition can be pushed, is the modern methodological reference this section’s use of the same partition follows [45].

Table 3 collects what the neutral share process predicts for the five families of statistics an empirical paper on market selection actually reports — size distributions, growth rates, concentration, selection terms, and reallocation work — with the source of each prediction. None of the six rows requires a selection mechanism to hold.

The point of assembling the table is not economy of exposition; it is that a single pair of

numbers, σ_ε and HHI, both already sitting in any panel of firm shares and growth rates, pins down every row at once. An analyst who has fit the growth-rate dispersion and computed the Herfindahl index has, without further assumption, already computed the neutral prediction for the cohort's size distribution, its concentration path, the spread its selection term should show if nothing were being selected, and the churn floor its reallocation statistic must clear. Before crediting any departure from that prediction to fitter firms winning, Proposition 2 says: compute what the statistic would be at the observed $(\sigma_\varepsilon, \text{HHI})$ under pure noise, and ask whether the data sit outside that range. §5 builds the test that answers this question exactly, without needing to know the shape of ε 's distribution at all.

5 The test

Section 4 ended on a specific instruction: before crediting any departure from the neutral share process to fitter firms winning, compute what the reported statistic would be at the observed $(\sigma_\varepsilon, \text{HHI})$ under pure noise, and ask whether the data sit outside that range. Carrying the instruction out for the Price selection term Sel needs a sampling distribution for Sel under H_0 , and here the tent-shaped growth-rate evidence of §3 forecloses the easy route. An asymptotic normal or t -reference distribution for Sel or $\hat{\beta}$ assumes the growth shocks ε_i are well behaved enough, and numerous enough per industry, for a central-limit argument to have taken hold; Stanley et al.'s Compustat evidence and Bottazzi and Secchi's fit of the Subbotin family to fifty-five Italian sectors both say the shocks are Laplace-shaped, not Gaussian, with fourth moments heavier than a normal reference tolerates [9, 103], and a typical industry-period panel has, at best, a few hundred firms — not the asymptotic regime any such approximation needs to be trusted. What is needed is a test whose null distribution is exact for whatever the data's own shocks look like, at whatever N the industry happens to have. Fisher supplied the principle six decades before market-share panels existed, in the same design-of-experiments argument that gave randomization inference its founding example: fix everything the null allows to vary, randomize only the assignment the null claims does not matter, and the resulting distribution of any statistic computed over that randomization *is* the null distribution, exactly, regardless of what the underlying shocks look like [42]. Proposition 3 applies the principle to Sel directly.

Proposition 3 (The permutation test). *Let the data be a panel $(s_i(t), s_i(t+1), z_i)$ for one industry-period. Define the statistics*

$$T_{\text{Sel}} = \frac{\text{Sel}}{\hat{\sigma}_{\text{Sel}}}, \quad T_{\mathcal{W}} = \frac{\mathcal{W} - \frac{\hat{\sigma}_\varepsilon^2}{2}(1 - \text{HHI})}{\hat{\sigma}_{\mathcal{W}}}, \quad \hat{\beta} \text{ from } \ln w_i = a + \hat{\beta}(z_i - \bar{z}) + e_i.$$

The null hypothesis H_0 (neutrality) is that w is independent of z conditional on s . Under H_0 the joint distribution of $(w_1, \dots, w_N)$ is invariant to any permutation of the labels of z across firms; hence the exact null distribution of any statistic that depends on the data only through the pairing of z with (s, w) — Sel, $\hat{\beta}$, and T_{Sel} among them — is obtained by recomputing the statistic over random permutations of z , and

$$p_{\text{perm}} = \frac{1 + \#\{\text{permutations with statistic} \geq \text{observed}\}}{1 + \#\text{permutations}}$$

is an exact p -value [38, 42, 73].

Ernst’s survey states the exactness argument in its most general form: whenever a null hypothesis implies a specific symmetry — here, exchangeability of the z -labels across firms — any statistic’s distribution over the group of relabelings that respects that symmetry *is* its conditional null distribution, with no appeal to sample size or to the shape of the underlying shocks, and Good’s manual turns the same principle into the standard applied toolkit [38, 51]. Three properties follow that make the test more than a computational convenience. First, it is distribution-free: nothing in Proposition 3 assumes ε_i is Gaussian, so the Laplace tails documented in §3 cost the test nothing — the permutation distribution is built from the data’s own realized shocks, heavy tails included, rather than referred to a reference family the shocks might not obey. Second, the test conditions on the observed shares $s_i(t)$: because Sel’s null variance is $\sigma_\varepsilon^2 \sum_i s_i^2 (z_i - \bar{z})^2$ (Proposition 2(i)), a concentrated industry’s permutation distribution is mechanically wider than a fragmented one’s at the same σ_ε , so $\sqrt{\text{HHI}}$ enters the test’s own null exactly where Proposition 1 says a market’s dispersion should enter it, without the analyst choosing a weighting scheme by hand. Third, T_W is *not* testable this way, and the paper states the limitation rather than gliding past it: W does not depend on the pairing of z with (s, w) that a relabeling permutes, so its permutation distribution carries no information at all about the null; its reference distribution has to be simulated or bootstrapped over $\hat{\sigma}_\varepsilon$ instead, in the manner of Efron and Tibshirani and of Davison and Hinkley [31, 37] — the price of a statistic, unlike Sel, that is a function of s and w alone and touches z nowhere.

Turning the proposition into a reporting protocol takes five steps.

1. Compute the observed statistic — Sel, $\hat{\beta}$, or T_{Sel} — on the panel exactly as recorded.
2. Permute the labels of z across firms, holding $s_i(t)$, $s_i(t+1)$, and hence w_i , fixed, 10,000 times.
3. Recompute the statistic on each of the 10,000 relabeled panels.
4. Report p_{perm} from the formula of Proposition 3, and express the observed statistic in units of the permutation distribution’s own standard deviation.
5. Repeat the first four steps industry by industry, and combine the resulting p -values across industries by a Fisher or Stouffer method.

The fifth step matters as much as the first four: a single industry’s permutation test cannot, on its own, distinguish one unlucky reshuffling from an economy-wide selective pressure, and combining across industries is the only way to accumulate evidence a one-industry test structurally cannot supply on its own.

Ten thousand permutations is a Monte Carlo approximation to the full permutation distribution, not the exact enumeration Fisher’s original argument imagines: for an industry of $N = 200$ firms the number of distinct relabelings, $200!$, is far beyond anything a computation could enumerate, so p_{perm} as reported by the five-step protocol is itself estimated with a controllable Monte Carlo error that shrinks as more permutations are drawn — a distinction Good’s manual makes explicit [51] and that this paper’s own simulations respect, using 10,000 reshufflings in the general protocol and 4,000 in the illustrative run of Figure 2 below. The distinction matters for how close to the boundary a reported p_{perm} can be trusted, but not for the exactness argument itself: the Monte Carlo estimate converges to the same conditional null distribution Proposition 3 defines, regardless of how many permutations are affordable to draw.

Running the five steps is also not the only route to the same question. Scoring a neutral and a selective specification against each other by an information criterion, in the manner of Burnham and Anderson, treats “selection or drift” as a model-comparison problem rather than a significance test, and is the natural complement to the p -value this protocol reports rather than a substitute for it [11].

The kinship with §2’s tests is exact rather than merely thematic, and Appendix B states each pairing in full: like the agreement Ewens’s sampling formula and Watterson’s estimator strike under strict neutrality, T_{Sel} checks an observed statistic against the one parameter pair, $(\sigma_\varepsilon, \text{HHI})$, that should have generated it under pure drift; like Tajima’s D , it standardizes a realized quantity by its own null spread rather than by an assumed reference distribution; like the Hudson–Kreitman–Aguadé test, its fifth step turns single-industry p -values into one joint statement across industries; and like the McDonald–Kreitman test, it needs no assumption about population size beyond what the conditioning itself supplies.

Figure 2 shows the procedure run once, on a simulated industry of $N = 200$ firms (Appendix C gives the full setup). Under neutrality the observed selection term is $\text{Sel} = -0.0046$, and its permutation distribution, built from 4,000 reshufflings of the productivity labels, gives $T_{\text{Sel}} = -0.91$ and $p_{\text{perm}} = 0.82$: correctly not rejected. With selection intensity $\beta = 0.10$ added to the same industry, the observed term is $\text{Sel} = +0.0165$ — positive, and larger than most of what chance produces — yet $T_{\text{Sel}} = 1.26$ and $p_{\text{perm}} = 0.10$: not rejected, because one industry of two hundred firms with growth noise of this size does not carry enough information to see a ten-percent selection intensity, exactly as the power surface of §6 says. With $\beta = 0.20$ the observed term is $\text{Sel} = +0.0161$, $T_{\text{Sel}} = 2.94$ and $p_{\text{perm}} = 0.002$: rejected. The two selective runs are the entire justification for the test: an observed Sel of 0.0165 was not evidence of selection and an observed Sel of 0.0161 was, because the pairings chance could have produced just as easily differed in spread between the two — a positive Sel carries no verdict until the pairing that produced it has been checked against them.

A test that rejects H_0 correctly at $\beta = 0.10$ is only useful if it would also reject at the smaller selection intensities the empirical reallocation literature actually reports, and on the smaller, more concentrated industries that literature actually studies; §6 turns to exactly that question.

6 Power, and the drift barrier of markets

A test that never rejects a true null is one kind of failure; a test that never rejects a false one is another, and it is the specific worry §5 already raises for a permutation procedure run on an industry-period panel of a few hundred firms. Exactness, in the sense of Proposition 3, says nothing about how often the test finds selection when selection is actually there — an exact test can still be a weak one, correctly sized and almost never rejecting. This section reports how the permutation test’s power behaves as the parameters of an industry-period panel change, and what that power surface implies about how much productivity difference a market can even see.

Proposition 4 (Power). *Against the alternative $w_i = 1 + \beta(z_i - \bar{z}) + \varepsilon_i$ with Laplace ε , the power of the permutation T_{Sel} test at size 0.05 rises with N , with σ_z , and with β , and falls with HHI at fixed N ; Appendix C reports the power surface for $N \in \{50, 200, 1000\}$, $\sigma_\varepsilon \in \{0.10, 0.30\}$, $\beta \in \{0, 0.02, 0.05, 0.10, 0.20\}$, $\sigma_z = 0.30$, and two share structures (equal shares; Zipf-like with*

HHI ≈ 0.05).

Corollary 3 (The drift barrier of markets). *A productivity difference δz is invisible to market selection when $\beta \delta z$ is small against the drift scale $\sigma_\varepsilon \sqrt{\text{HHI}}$ — concentrated industries have a higher drift barrier and tolerate larger inefficiency differences without selecting on them; the minimum effective difference scales as*

$$\delta z_{\min} \propto \frac{\sigma_\varepsilon \sqrt{\text{HHI}}}{\beta}.$$

The corollary transcribes, rather than newly derives, a result population genetics already owns. Ohta’s nearly neutral theory and Lynch’s later formalization of the same criterion — fixation governed by $S = 2N_g s$, so that an effect too small relative to $1/N_g$ is driven by drift regardless of its sign [80, 92] — is exactly Corollary 3 once N_g is read as $N_e = 1/\text{HHI}$ and s as $\beta \delta z$: a concentrated market is a small- N_e population in Ohta and Lynch’s own sense, and a productivity gap that would be decisive in a fragmented industry is invisible in a concentrated one running at the same growth-shock variance and the same selection strength.

The corollary’s mechanism is not a separate assumption; it falls out of the parameter this paper has used throughout. $N_e = 1/\text{HHI}$ is, by Proposition 1, the number of firms a market’s share distribution behaves as if it had; a permutation test’s power to detect a covariance between z and w depends, in the ordinary way of any covariance-type test, on how many effectively independent comparisons the data supply, and a concentrated industry supplies fewer of them than its raw firm count N suggests, because a handful of large firms carry almost all of the share-weighted variance Sel’s null depends on. Table 5’s Zipf-like columns hold N fixed at every row and still detect less than the equal-share columns: at $N = 1,000$ and $\sigma_\varepsilon = 0.30$ the gap is 0.95 against 0.59 for $\beta = 0.10$, a difference attributable to nothing but HHI, since every other parameter is held fixed. Concentration does two things to an industry at once, and Proposition 1 and Proposition 4 between them state both: it drifts upward on its own under pure noise, and it lowers the effective sample size available to notice, should a firm actually be better than its rivals.

The growth-shock variance σ_ε plays no role in Proposition 4’s own list of what raises power, because it is a nuisance parameter for this test rather than a signal one: it is the scale of the very drift the permutation is built to condition out, and a market with noisier neutral growth needs a correspondingly larger N or β to reach the same power against it as a quieter one would. This is why §3’s Laplace-tailed, sizeable growth-rate dispersion is not merely a technical reason to prefer a permutation test over an asymptotic one, as §5 argued; it is also the reason the resulting power surface, evaluated at that dispersion, looks as unforgiving as it does.

An underpowered test carries a specific interpretive risk that exactness does nothing to remove: failing to reject H_0 is evidence for neutrality only in proportion to the test’s power against whatever alternative is actually operating, and a reading practice that treats $p > 0.05$ as though it settled the question inherits every one of Table 5’s low-power cells as a false confirmation of drift. This is the mirror image of the risk §4 raised for a rising Herfindahl index: there, a neutral market could be mistaken for a selective one; here, at $\beta = 0.10$, $N = 200$ and $\sigma_\varepsilon = 0.30$ under equal shares, the test’s own power of 0.33 means a genuinely selective market at that intensity is missed roughly twice as often as it is caught. Proposition 4 does not remove this risk — no amount of exactness in a test’s size compensates for weak power against the alternative that matters — but it makes the risk computable rather than merely suspected.

Appendix C’s power surface (Table 5, Figure 3) is where Proposition 4’s qualitative claims become numbers, and the numbers are less comfortable than the qualitative statement alone suggests. At the growth-shock variance the Compustat and Italian evidence of §3 actually reports, $\sigma_\varepsilon = 0.30$, an industry-period panel of $N = 200$ firms detects a selection intensity of $\beta = 0.10$ only 0.40 of the time with equal shares and 0.28 of the time with Zipf-like shares, and detects $\beta = 0.05$ only 0.14 of the time under either — a test that misses a five-percent-per-period selection intensity roughly six times out of seven. Power recovers only once N reaches four figures: at $N = 1,000$ the test detects $\beta = 0.10$ with probability 0.93 under equal shares, but only 0.56 under Zipf-like shares, the shortfall being exactly Corollary 3’s statement that concentration — which lowers N_e — raises the drift barrier at fixed N and σ_ε . None of this is a defect in the test: it holds its nominal size correctly in every cell Appendix C reports. It is a statement about how much information a few hundred firms’ worth of noisy growth rates actually carries about a selection intensity of realistic size.

That comparison matters because realistic size is not a hypothetical. Paper No. 1’s selection-accounting exercise finds settings — Bartelsman, Haltiwanger and Scarpetta’s cross-country Olley–Pakes comparison among them, with a labor-productivity covariance term near 50 log points in the United States against 20 to 30 in Western Europe — consistent with a selection gradient well above the $\beta = 0.10$ this section’s power surface treats as its largest tabulated case [7, 30]. It is the more modest coefficients reported elsewhere in the productivity–growth literature, on the order of $\beta = 0.02$ to 0.05 once translated into this section’s units, that Table 5 says a single industry of a few hundred firms, at the growth-noise level real panels display, is unlikely to detect at conventional power even when the underlying selection gradient is genuinely nonzero. A weak or statistically fragile productivity–growth coefficient reported from a study of that size is therefore compatible with two very different worlds — one in which selection is genuinely weak, and one in which it is not weak at all but the design could not have found it — and Proposition 4 is what lets the two be told apart, once a study’s N , HHI and growth-rate dispersion are on the table alongside its coefficient.

None of this argues against measuring selection intensity in market-share panels; it argues for reporting, alongside any estimated β or Sel, the three numbers that Proposition 4 says determine whether the estimate could have been anything else — N , HHI, and σ_ε — so that a reader can locate the study on a surface like Table 5 rather than take the point estimate on faith. Paper No. 2’s replicator-dynamics account of within-market competition assumes a selection gradient acting continuously on every firm in a population [26]; Table 5 is the empirical answer to how large that gradient, and how large the industry, would have to be before a market-share panel could confirm the assumption rather than merely fail to contradict it. §7 puts exactly those numbers, where the published record supplies them, next to the coefficients already in print.

7 Re-reading the evidence

The discipline §6 sets up is simple to state and easy to apply unevenly, so it is stated once, plainly, before applying it: for each finding below, report (a) the result as its authors published it, (b) what the null of §4 predicts for that statistic at the study’s approximate sample size and concentration, and (c) whether the study’s own design could, at conventional power, have told the reported effect apart from drift. No study below is accused of getting anything wrong. Several of the studies below say, in their own words, that the effect they find is weak; the

question this section adds is not whether they are honest about that weakness but whether their design was capable of finding a strong effect had one been there. Where (b) or (c) cannot be computed because the published paper does not report the sample size or the dispersion Table 5 needs, that is stated too, rather than filled in by assumption.

This reallocation-accounting tradition did not start with the seven findings below. Baily, Hulten and Campbell’s early plant-level study documented large dispersion and continual churning of productivity across individual US manufacturing plants, an empirical foundation the later decomposition literature is built on [6], and Foster, Haltiwanger and Krizan’s canonical decomposition showed that reallocation of outputs and inputs across establishments, together with net entry, contributes substantially to aggregate productivity growth [43] — the accounting convention every finding below either extends or is read against. Neither study reports a firm count and growth-rate dispersion at the resolution Table 5 needs, so neither is added as its own row; both are cited here as the lineage the seven studies below descend from.

Bottazzi, Dosi, Jacoby, Secchi and Tamagni’s comparison of Italian and French industrial panels reports that “heterogeneity in efficiencies primarily yield persistent profitability differentials, whereas the relationships of corporate growth with either productivity or profitability appear much weaker, if at all existent”; in their own worked example, productivity accounts for roughly one percent, and typically under five percent, of the variance in growth rates [10]. Neither a firm count nor a growth-rate dispersion at the resolution Table 5 needs survives into the verified record, so neither the implied power band nor a verdict distinguishing weak selection from drift can be computed; the finding is reported faithfully as what it is, a small explained-variance share from a panel of unstated size.

Dosi, Moschella, Pugliese and Tamagni’s four-country comparison finds a somewhat larger, still modest productivity contribution — a median elasticity near 0.2 and a productivity share of explained growth variance of 14 to 19 percent across France, Germany, the UK and the USA [33], larger than Bottazzi et al.’s own estimate but modest on the paper’s own reading. No sample size or dispersion figure survives into the verified extraction here either, so the same undetermined verdict applies for the same reason.

Coad’s test of “growth of the fitter” on French manufacturing firms is the one entry in this section with an approximate sample size, though the figure is a secondary-source paraphrase rather than one confirmed from the primary text, which resisted every access route attempted: on the order of 8,400 firms tracked from 1996 to 2004, non-parametric plots showing no obvious profit–growth relationship and parametric estimates finding, at best, a small and sometimes statistically insignificant positive association [18]. An N in the thousands sits an order of magnitude above the largest size Table 5 tabulates, where power against $\beta = 0.05$ already reaches 0.46 under equal shares at $N = 1,000$; by Proposition 4’s monotonicity in N , the implied power at $N \approx 8,400$ is higher still for any of the table’s tabulated selection intensities. A weak profit–growth link found at that scale is not obviously a power failure, and the more natural verdict is drift-consistent.

Foster, Haltiwanger and Syverson’s decomposition of physical and revenue productivity is the sharpest entry in this section, because it reports both a sample size and a dispersion: 17,699 plant-year observations and a within-industry standard deviation of physical TFP of 0.26 log points, close to this paper’s own $\sigma_z = 0.30$ [44]. A one-standard-deviation increase in physical productivity is associated with a 1.2 percent decline in exit probability, against a 5.0 percent decline associated with a one-standard-deviation demand shock — the productivity effect is three to four times smaller than the demand effect, on the authors’ own comparison.

At $N = 17,699$, again an order of magnitude above the table’s ceiling, Proposition 4’s rise in power with N makes it unlikely that the small productivity effect is a false negative the design was too weak to find; the verdict is selection, smaller than the demand-driven component the same study isolates but not plausibly drift alone.

Bartelsman, Haltiwanger and Scarpetta’s cross-country Olley–Pakes covariance comparison — about 50 log points within United States manufacturing industries against 20 to 30 in Western Europe [7] — is not, in the strict sense Proposition 3 needs, a T_{Sel} -type statistic at all: the OP covariance is a level computed from the allocation of activity shares, not a permutation-tested regression coefficient, and no firm count or HHI accompanies it in the verified record. Table 5’s power surface is built for T_{Sel} under this paper’s specific null and cannot be mapped onto an OP covariance without assumptions the source does not supply; the verdict is undetermined by construction of the comparison, not because a thirty-log-point US–Europe gap looks like a plausible chance event.

Hsieh and Klenow’s cross-country accounting exercise supplies this comparison’s natural benchmark: reallocating capital and labor to equalize marginal products to the extent already observed within the United States would raise measured manufacturing TFP by 30 to 50 percent in China and 40 to 60 percent in India [57]. Like the Bartelsman–Haltiwanger–Scarpetta comparison just above, this is a level computed from the allocation of activity rather than a permutation-tested T_{Sel} coefficient, and it carries no firm count or HHI that would let Table 5 speak to it directly; it is reported here as the scale of the reallocation gap the accounting tradition treats as its headline number, not as an eighth row in Table 4.

Decker, Haltiwanger, Jarmin and Miranda’s finding that employment’s marginal responsiveness to productivity shocks fell to roughly half its 1990s peak in high-tech manufacturing and two-thirds economy-wide, while the underlying dispersion of shocks did not decline [32], is drawn from the full universe of US Census Bureau business microdata, a scale far beyond anything Table 5 tabulates, though the extraction reports no single firm count or growth-rate dispersion at the industry-period level needed to place the finding on the table directly. At that scale a decline of the reported magnitude is unlikely to be a sampling artifact; the paper’s central claim is a weakening selection gradient, not a level of $\text{Var}_s(z)$ that has fallen, and the verdict is selection, specifically a declining one.

Cantner and Krüger’s decomposition of German manufacturing productivity growth across eleven two-digit industries between 1981 and 1998 attributes the large post-reunification productivity gains mainly to structural change between industries rather than within-firm learning [13]. The decomposition is computed at the industry level, and no firm-level N or dispersion figure is available in the verified record to place the finding on Table 5 at all; here, uniquely in this section, even the statistic’s unit of analysis puts a power reading out of reach, and the verdict is undetermined.

Two of the seven findings above read as selection once their scale is taken into account, one reads as drift-consistent once its scale is taken into account, and four cannot be placed on the power surface at all because the published record does not report what Table 5 needs — a firm count, a growth-rate dispersion, or a statistic of the right type. That four-out-of-seven majority is itself a finding, not a gap in this section’s search: it is the direct evidence that most of the published productivity–growth literature was not designed to answer the question this paper asks, and reports coefficients without the companion numbers that would let a reader compute whether the design could have answered it.

The model for this kind of discipline is not economic; it is McGill’s test of Hubbell’s zero-

Study	N	Reported effect	Implied power (Table 5)	Verdict
Bottazzi et al. 2010 [10]	not reported	productivity $\approx$ 1–5% of growth variance	not computable	undetermined
Dosi et al. 2015 [33]	not reported	productivity $\approx$ 14–19% of growth variance	not computable	undetermined
Coad 2007 [18]	$\approx 8,400^\dagger$	weak/inconsistent profit→growth link	$N \gg 1,000$; power above the table’s ceiling row	drift-consistent
Foster–Haltiwanger–Syverson 2008 [44]	17,699	1 SD TFPQ $\rightarrow$ 1.2% lower exit prob. (vs. 5.0% for demand)	$N \gg 1,000$; power above the table’s ceiling row	selection (small)
Bartelsman–Haltiwanger–Scarpetta 2013 [7]	not reported	OP covariance: US $\approx$ 50, Western Europe 20–30 log pts	statistic not on the T_{Sel} scale	undetermined
Decker et al. 2020 [32]	Census-scale, not itemized	responsiveness fell to 1/2–2/3 of 1990s peak	scale $\gg 1,000$; per-period β not reported	selection (declining)
Cantner–Krüger 2008 [13]	not applicable (industry-level)	post-reunification gains mainly structural change	not applicable (industry-level statistic)	undetermined

Table 4: Re-reading seven verified productivity–growth findings against the power surface of Table 5. † Coad’s sample size is a secondary-source figure, not confirmed from the primary text (see discussion above). “Not computable” and “not applicable” mark cases where the published record omits the sample size, dispersion, or statistic type Table 5’s power surface requires — not a judgment that the underlying study is flawed.

sum multinomial model [59] against the plain lognormal on the Barro Colorado Island tree census and the North American Breeding Bird Survey data (§2): a null taken seriously enough to be fit, quantified, and beaten on its own terms, reported without hedging — “the ZSM fail[s] to fit empirical data better than the lognormal distribution 95% of the time” [85] — rather than defended by appeal to the mechanism’s biological plausibility. Volkov, Banavar, Hubbell and Maritan had already supplied the testable distribution the comparison needed [109], and Alonso, Etienne and McKane’s reply supplies the discipline’s other half: a failed goodness-of-fit test rejects a specific prediction, not the whole research programme, because “neutral theory is also a stochastic theory, a sampling theory and a dispersal-limited theory” whose components can be probed one at a time [1]. Economics now has, in §4–§6, the null, the exact test, and the power calculation that ecology built and then used against itself; what the seven studies above show it has mostly lacked is not honesty but the habit of reporting the N , the dispersion, and the power alongside the coefficient — not because any of the seven got anything wrong, but because none was asked the question this paper is built to ask.

8 The drift barrier for markets

Section 2 stated Lynch’s drift-barrier principle in its own terms: fixation of a variant with selection coefficient s in a population of effective size N_g is governed by $S = 2N_g s$, “twice the ratio of the power of selection to the power of random genetic drift,” so that “if $1/N_g \gg |s|$, the population will be driven to a state expected under mutation pressure alone” [80]. Ohta’s nearly neutral theory supplies the same floor from the opposite direction: a mutation that is genuinely deleterious, not exactly neutral, is still fixed by drift once the population is small enough, so that evolution runs fastest precisely where selection is weakest relative to chance [92]. Charlesworth’s synthesis states what stands behind both results in one quantity: effective population size fixes simultaneously the level of neutral variation a population carries and “the effectiveness of selection relative to drift” [15]. None of this is a claim about how strong selection is in an absolute sense; it is a claim about a ratio, and a ratio has a denominator that can shrink as easily as a numerator can. Table 1 and Appendix B already write the translation of that ratio into markets in a single line. This section works out what the line costs once it is put to work on a policy problem.

Corollary 3 of §6 is the transcription this section takes up as a principle: a productivity difference δz is invisible to market selection when $\beta \delta z$ is small against the drift scale $\sigma_\varepsilon \sqrt{\text{HHI}}$, so that the minimum effective difference scales as $\delta z_{\min} \propto \sigma_\varepsilon \sqrt{\text{HHI}}/\beta$; concentrated industries have a higher drift barrier and tolerate larger inefficiency differences without selecting on them than fragmented industries at the same σ_ε and β .

The translation goes through $N_e = 1/\text{HHI}$, the same effective population size that Proposition 2 of §4 already uses to write Papers No. 1, 3 and 7 in one notation [30]. A market with HHI near zero has N_e large and a low barrier: almost any consistent productivity edge, however small, eventually shows up as a share gain, because the denominator of Lynch’s ratio is large. A market with HHI near one has N_e near one and a barrier close to its ceiling: the few firms left are, from selection’s point of view, nearly a founder population, and Corollary 3 says a substantial δz can sit there indefinitely without the market ever acting on it. Because Proposition 1 also shows N_e falling in expectation under pure noise, the barrier is not a fixed feature of an industry’s technology; it rises on its own, mechanically, exactly when concentration drifts upward for reasons that have nothing to do with fitness [30].

The barrier is not only a fact about what the market does; it is, in exactly the same units, a fact about what an outside observer can learn. Proposition 4’s power surface shows the permutation test T_{Sel} of §5 rising in power with N , with σ_z , and with β , and falling with HHI at fixed N (Figure 3)—not a separate econometric coincidence but the same ratio read twice. The statistic Sel and the test statistic built from it are both, to first order, weighted sums of the same growth shocks ε_i that generate δHHI in Proposition 1; a δz too small to move the market’s own reallocation by more than $\sigma_\varepsilon \sqrt{\text{HHI}}$ is too small to move $\hat{\beta}$ ’s sampling distribution by more than the same amount, for the same reason. A concentrated industry is therefore doubly opaque: selection cannot see a small δz inside it, and neither, at any realistic panel length, can the econometrician who goes looking for the same effect with the tools of §5. §7’s study-by-study re-examination returns to this identity directly: a reported “weak selection” coefficient in a concentrated industry is, until the drift floor is subtracted, indistinguishable from a well-powered test correctly finding nothing and an underpowered test failing to find something real.

That identity has a policy reading that the antitrust literature has not stated in these terms.

Hypothesis 1 (antitrust as effective-population management). Because $N_e = 1/\text{HHI}$ sets both the rate at which concentration self-reinforces (Proposition 1) and the location of the drift barrier (Corollary 3), a merger that raises HHI is, independent of any price effect the standard unilateral-effects analysis prices, a policy-induced reduction in the market’s effective population size—and therefore a policy-induced increase in $\delta z_{\min}$, the inefficiency gap the market will subsequently tolerate without selecting on it. On this reading, merger review already regulates N_e whether or not it says so. *Estimation design:* for a cohort of reviewed mergers, estimate $\hat{\sigma}_\varepsilon$ from the Subbotin/Laplace fit to each industry’s growth-rate distribution [9] and $\hat{\beta}$ from the Olley–Pakes covariance between market share and productivity [94], computed on a dynamic decomposition that separates entering and exiting units from incumbents [86]; form the drift-barrier proxy $\hat{\sigma}_\varepsilon \sqrt{\text{HHI}}/\hat{\beta}$ before and after consummation, and compare its change across consummated mergers against a control group of mergers reviewed in the same period but blocked or abandoned, to net out the selection of which deals get proposed at all. *Refutation condition:* no divergence between the two groups’ post-review drift-barrier proxy would say the HHI that antitrust screens on is not, in fact, moving the quantity Corollary 3 says decides what the market can subsequently see.

Hypothesis 2 (merger review as a change in the drift scale). A merger’s effect on future selection is fully summarized, under Corollary 3, by its effect on the ratio $\sigma_\varepsilon \sqrt{\text{HHI}}/\beta$, not by its effect on HHI alone: two mergers producing identical post-merger HHI can leave very different drift barriers behind if one of them also removes a rival whose presence had been raising β (a competitor uniquely quick to reward productivity gains with share) or lowering σ_ε (a source of unusually predictable rather than noisy rivalry). *Estimation design:* run the permutation test of Proposition 3 on the same industry’s panel before and after the merger—or on the merging industry against a synthetic-control industry matched on pre-merger HHI, $\hat{\sigma}_\varepsilon$ and $\hat{\beta}$ —and read the change in its power against a fixed effect size, taken from Proposition 4’s power surface (Appendix C), as the direct measure of how far the merger moved the drift scale, holding the effect the antitrust agency is trying to protect fixed. *Refutation condition:* a merger that clears the standard structural-presumption HHI threshold but leaves this power estimate essentially unchanged would show that the HHI screen alone is a poor proxy for the drift-scale change the corollary identifies as the object of concern, and that $\hat{\sigma}_\varepsilon$ and $\hat{\beta}$ need to be estimated directly rather than proxied by concentration alone.

The barrier also reframes three results already assembled elsewhere in this series. Paper No. 3’s sorting-work statistic $\mathcal{W}$ has a neutral expectation of $\frac{\sigma_\varepsilon^2}{2}(1 - \text{HHI})$ (Proposition 2(ii)) built from exactly the same σ_ε and HHI that set $\delta z_{\min}$ here [29]. Below the barrier, whatever churn a concentrated market produces is, by construction, the neutral floor itself rather than evidence of firms being sorted by fitness: Schumpeter’s demon can appear to do costless sorting work in exactly the regime where Corollary 3 says no productivity difference is actually being selected on, and the two papers’ statistics are then reporting the same noise from two directions rather than confirming each other. Paper No. 4’s account of persistent productivity dispersion — a pattern Syverson’s survey of the productivity literature finds ubiquitous across businesses and industries generally, not peculiar to any one dataset [107] — attributes stuck firms to a rugged technology landscape—local peaks separated by valleys no incremental search crosses [28]. Corollary 3 supplies a second, market-structure route to the identical observable: a firm can sit below a rival on a landscape that is not rugged at all and still never be selected against, if the gap between them is smaller than $\sigma_\varepsilon \sqrt{\text{HHI}}/\beta$: a concentrated market flattens the selective gradient the same way a rugged landscape flattens the fitness gradient, so a ruggedness

estimator fit without conditioning on HHI risks crediting the landscape with dispersion the drift barrier alone would have produced on a smooth one. Paper No. 9’s two-mechanism account of profit-persistence decay—rivals imitating an advantage away or answering it with an arms race—is identified, in its own framework, from the joint behavior of profit autocorrelation and contested-input expenditure [25]. An arms race requires a rival to detect the gap it is racing to close; Corollary 3 says that detection itself has a floor, so a concentrated industry showing slow profit-persistence decay may be showing a race that markets cannot see well enough to run, rather than a coordinated truce, and Paper No. 9’s identification strategy is sharpest exactly where this paper’s barrier is lowest.

None of these three connections asks Papers 3, 4 or 9 to be re-derived; each asks that their statistics be read alongside $N_e = 1/\text{HHI}$ before a null result is credited to the mechanism each paper proposes rather than to the floor this section states. §7 takes up that reading study by study, and §9’s reporting standard is the natural place to add a fourth number, alongside p_{perm} , N_e , and power at the reported effect: the drift-barrier ratio itself, so that a claimed selection intensity is always reported next to the floor below which no test, and no market, could have told it from noise.

9 Conclusion and program

The preceding sections state one null, one exact test, and one power surface. None of the three is, by itself, a verdict on any industry; together they are a standard a selection claim made on market-share data can be held to, in the same way Kimura’s neutral theory and its tests became a standard a selection claim made on sequence data has been held to since 1968 [68]. What changes once the standard exists is not any particular finding of this series’ companion papers, but what each of their central objects now owes a reader before it is read as a discovery about selection rather than a description of noise.

1. The Price selection term of **(author?)** [30] now has an exact null variance (Proposition 2i) and a permutation p -value it did not have before: a measured Sel is reportable as significant or not, not merely as positive.
2. The claim that market competition performs an optimizing computation [26] now has a floor: the neutral share process performs the identical random walk on the simplex with no computation being done at all, so a reallocation must be shown to exceed what the walk alone produces before it is credited to competition.
3. The reallocation-work statistic of **(author?)** [29] now has the neutral benchmark $\frac{\sigma_\varepsilon^2}{2}(1 - \text{HHI})$ it lacked (Proposition 2ii), together with Proposition 3’s caveat that $\mathcal{W}$ is not permutation-testable, which directs that paper to a bootstrap over ε rather than a shuffle of labels.
4. A rugged fitness landscape with multiple stable configurations [28] is also what a drift walk on the simplex produces with no landscape at all; a claimed peak must be shown to survive reshuffling the productivity labels that are supposed to define it.
5. A rising share for a cooperative organizational form [22] is not, by itself, evidence that cooperation is fitter: Proposition 1’s concentration-drift formula gives the rate of rise a skeptic is entitled to subtract first.

6. A routine’s persistence across periods [27] needs the same drift baseline behind it that any neutral marker needs: some routines persist and spread for the reason some alleles reach high frequency under neutrality, and the permutation test is exactly the tool for asking whether a routine’s frequency tracks performance or only time.
7. The lock-in number $N_e k$ (Proposition 2iii) gives (author?)’s claim that early accidents can freeze into permanent standards [23] its quantitative boundary: below that threshold, lock-in is drift and not a demonstrated advantage; above it, drift cannot be blamed.
8. A currency’s rising share of transaction volume [24] is concentration like any other, and Proposition 1 says part of that rise is guaranteed by chance alone; a claim that money selects itself for useful properties owes readers the fraction of the rise the null does not explain.
9. Sustained turnover at the top of a competitive ranking, the signature a Red Queen dynamic predicts [25], is also what the null’s concentration drift already supplies at the same N_e ; the permutation test is what tells the two rates of turnover apart.

None of these nine clauses overturns its paper’s argument. Each narrows it to what the argument can show once the null it was implicitly competing against is written down and given a p -value.

The empirical program

The propositions of §4–§6 are not, by themselves, an empirical program; they are a null and a test. The program is what a reader does with them before believing a selection claim made on market-share data. First, compute the panel’s Herfindahl index and the growth-shock scale σ_ε from the data at hand, and apply the drift-scale corollary of Proposition 4 as a screen before running anything: given the panel’s N and HHI, what selection intensity β would the design need to detect at reasonable power, and does the effect size under discussion clear that bar (Table 5)? A panel that fails the screen — a few hundred firms, growth noise at the Laplace scale documented in §3, and a concentrated industry — should report the drift floor and stop, rather than run a test with little chance of separating signal from noise and report an uninformative result as though it were evidence of no selection. Second, where the screen is cleared, run the five-step permutation protocol of Proposition 3 — Fisher’s randomization logic [42], mechanized — on public longitudinal panels: Census-style business registers and the national statistical agencies’ establishment panels the companion papers already draw numbers from. Report p_{perm} , $N_e = 1/\text{HHI}$, and the achieved power at the estimated β together, every time, as one reporting unit rather than a bare significance star; combine industries by Fisher’s or Stouffer’s method rather than counting how many crossed $p < 0.05$ on their own. This is the reporting standard the introduction promised: every selection claim in market-share data states its drift floor before it states its finding.

What would falsify this paper’s claims

Four things would. If real panels drawn from plausibly neutral sectors show Herfindahl paths that systematically diverge from Proposition 1’s formula, the assembled null is wrong, not merely imprecise. If the permutation test fails to hold its nominal size on real data — reject

neutrality far more often than 0.05 where selection is independently implausible — then the exchangeability condition Proposition 3 rests on is violated in practice, most likely by growth shocks correlated across firms within a period rather than the independent shocks the Definition assumes, and the test’s exactness claim does not survive contact with real panels. If well-powered re-analyses — large N , low concentration, the top rows of Table 5 — of the studies re-read in §7 continue to find only weak selection, then this paper’s central empirical wager, that some of the literature’s weak selection findings are underpowered rather than genuinely weak, is wrong, and the honest conclusion becomes that selection in those markets really is weak. And if the productivity measures on which any of this is run turn out, as Foster, Haltiwanger and Syverson’s physical-versus-revenue distinction warns, to be measuring market power rather than efficiency [44], every number in this program is a statement about price-setting ability dressed as a statement about fitness, and the re-reading of §7 inherits the same warning rather than resolving it.

Closing

This paper does not show that any industry is or is not selecting on the efficiency of its firms. It shows what a reader must compute before answering that question for any industry, and it reports, without dressing the finding up, that when the computation is done at the sample sizes and concentration levels real industry panels actually have, the honest answer is very often that the data cannot yet tell. That is a smaller claim than a discovery of pervasive market selection would be, and a smaller claim than a debunking of the selection-accounting literature would be. Population genetics did not resolve its own selectionist–neutralist argument by building the neutral theory and its tests; it made the argument answerable, decade after decade, on the same shared instrument, willing on at least one documented occasion to watch its own null lose to the data outright [85]. Economics has now been offered the same instrument for its own oldest argument about whether markets reward the efficient. What it does with the offer is not this paper’s to decide.

A Proofs

A.1 The second-order expansion of concentration drift (Proposition 1)

Fix a period and write $s_i \equiv s_i(t)$, $w_i = 1 + \varepsilon_i$, with $\varepsilon_1, \dots, \varepsilon_N$ independent across i , mean 0, variance σ_ε^2 . Let $\bar{\varepsilon} \equiv \sum_i s_i \varepsilon_i$ be the share-weighted shock, so that the normalizing sum in Definition 1 is $\sum_j s_j w_j = 1 + \bar{\varepsilon}$. Since $s_i(t+1) = s_i w_i / (1 + \bar{\varepsilon})$,

$$\text{HHI}(t+1) = \sum_i s_i(t+1)^2 = \frac{\sum_i s_i^2 (1 + \varepsilon_i)^2}{(1 + \bar{\varepsilon})^2},$$

which is the exact object Proposition 1 expands to second order in ε .

Numerator and denominator, separately. Expanding the numerator,

$$\sum_i s_i^2 (1 + \varepsilon_i)^2 = \sum_i s_i^2 (1 + 2\varepsilon_i + \varepsilon_i^2) = \text{HHI} + \underbrace{2 \sum_i s_i^2 \varepsilon_i}_{=: A} + \underbrace{\sum_i s_i^2 \varepsilon_i^2}_{=: B},$$

with $A = O(\varepsilon)$ first order and $B = O(\varepsilon^2)$ second order. Expanding the denominator's reciprocal about $\bar{\varepsilon} = 0$ using $(1+x)^{-2} = 1 - 2x + 3x^2 - O(x^3)$ with $x = \bar{\varepsilon} = O(\varepsilon)$,

$$\frac{1}{(1+\bar{\varepsilon})^2} = 1 - 2\bar{\varepsilon} + 3\bar{\varepsilon}^2 + O(\varepsilon^3).$$

Multiplying and collecting orders. Multiplying the two expansions and retaining every term through second order in ε (terms such as $A\bar{\varepsilon}^2$ and $B\bar{\varepsilon}$ are third order and are dropped),

$$\text{HHI}(t+1) = \text{HHI} + \underbrace{(A - 2\text{HHI}\bar{\varepsilon})}_{\text{first order}} + \underbrace{(3\text{HHI}\bar{\varepsilon}^2 - 2A\bar{\varepsilon} + B)}_{\text{second order}} + O(\varepsilon^3).$$

Substituting $A = 2 \sum_i s_i^2 \varepsilon_i$, the second-order bracket is $3\text{HHI}\bar{\varepsilon}^2 - 4\bar{\varepsilon} \sum_i s_i^2 \varepsilon_i + \sum_i s_i^2 \varepsilon_i^2$. Because $\mathbb{E}[\varepsilon_i] = 0$ for every i , the first-order bracket has expectation zero:

$$\mathbb{E}[A - 2\text{HHI}\bar{\varepsilon}] = 2 \sum_i s_i^2 \mathbb{E}[\varepsilon_i] - 2\text{HHI} \sum_i s_i \mathbb{E}[\varepsilon_i] = 0,$$

so the leading contribution to $\mathbb{E}[\Delta\text{HHI}]$ is entirely second order, and it remains to take the expectation of the three pieces of the second-order bracket, using independence of $\varepsilon_1, \dots, \varepsilon_N$ throughout.

The three second-order expectations. First,

$$\mathbb{E}[3\text{HHI}\bar{\varepsilon}^2] = 3\text{HHI} \mathbb{E}\left[\left(\sum_i s_i \varepsilon_i\right)^2\right] = 3\text{HHI} \sum_i s_i^2 \text{Var}(\varepsilon_i) = 3\sigma_\varepsilon^2 \text{HHI}^2,$$

since $\mathbb{E}[\bar{\varepsilon}^2] = \text{Var}(\bar{\varepsilon}) = \sum_i s_i^2 \sigma_\varepsilon^2$ by independence (cross terms $\mathbb{E}[\varepsilon_i \varepsilon_j] = 0$ for $i \neq j$ vanish) and $\sum_i s_i^2 = \text{HHI}$. Second,

$$\mathbb{E}\left[-4\bar{\varepsilon} \sum_i s_i^2 \varepsilon_i\right] = -4 \sum_{i,j} s_j s_i^2 \mathbb{E}[\varepsilon_j \varepsilon_i] = -4 \sum_i s_i \cdot s_i^2 \cdot \sigma_\varepsilon^2 = -4\sigma_\varepsilon^2 \sum_i s_i^3,$$

because $\mathbb{E}[\varepsilon_j \varepsilon_i] = \sigma_\varepsilon^2$ when $i = j$ and 0 otherwise, leaving only the $i = j$ terms of the double sum. Third,

$$\mathbb{E}\left[\sum_i s_i^2 \varepsilon_i^2\right] = \sum_i s_i^2 \mathbb{E}[\varepsilon_i^2] = \sigma_\varepsilon^2 \sum_i s_i^2 = \sigma_\varepsilon^2 \text{HHI},$$

using $\mathbb{E}[\varepsilon_i^2] = \text{Var}(\varepsilon_i) + (\mathbb{E}[\varepsilon_i])^2 = \sigma_\varepsilon^2$. Summing the three, and adding back the zeroth-order carry-over term HHI from which ΔHHI is measured, gives the exact second-order expectation in the four pieces the proposition is built from — HHI itself, $+3\sigma_\varepsilon^2 \text{HHI}^2$, $-4\sigma_\varepsilon^2 \sum_i s_i^3$, and $+\sigma_\varepsilon^2 \text{HHI}$ — so that

$$\mathbb{E}[\text{HHI}(t+1)] = \text{HHI} + 3\sigma_\varepsilon^2 \text{HHI}^2 - 4\sigma_\varepsilon^2 \sum_i s_i^3 + \sigma_\varepsilon^2 \text{HHI} + O(\varepsilon^3),$$

and subtracting HHI from both sides,

$$\mathbb{E}[\Delta\text{HHI}] = \sigma_\varepsilon^2 \left(\text{HHI} - 4 \sum_i s_i^3 + 3\text{HHI}^2 \right) + O(\varepsilon^3),$$

which is Proposition 1. □

The equal-shares check. With N firms at $s_i = 1/N$, $\text{HHI} = \sum_i (1/N)^2 = 1/N$ and $\sum_i s_i^3 = N \cdot (1/N)^3 = 1/N^2$, so

$$\mathbb{E}[\Delta\text{HHI}] = \sigma_\varepsilon^2 \left(\frac{1}{N} - \frac{4}{N^2} + \frac{3}{N^2} \right) = \sigma_\varepsilon^2 \left(\frac{1}{N} - \frac{1}{N^2} \right) = \sigma_\varepsilon^2 \frac{N-1}{N^2},$$

strictly positive for every $N \geq 2$ and vanishing only in the degenerate case $N = 1$, a single firm holding the entire market — monopoly by definition, where there is nothing left to reallocate and $\Delta\text{HHI} \equiv 0$ identically, not merely in expectation.

The near-monopoly sign. Consider a dominant firm of share $1 - \delta$ for small $\delta > 0$, together with an arbitrary fringe of one or more firms whose shares sum to δ . Write $Q \equiv \sum_{\text{fringe}} s_j^2$ for the fringe's own contribution to HHI; because the fringe shares are each at most δ and sum to δ , $Q = c\delta^2$ for some $c \in (0, 1]$ depending on how the fringe is split ($c = 1$ if the fringe is a single firm; $c = 1/m$ if it is m equal firms), and the fringe's contribution to $\sum_i s_i^3$ is $R = \sum_{\text{fringe}} s_j^3 \leq \delta \cdot Q = O(\delta^3)$, negligible at the order kept below. Then, to $O(\delta^2)$,

$$\text{HHI} = (1 - \delta)^2 + Q = 1 - 2\delta + (1 + c)\delta^2, \quad \sum_i s_i^3 = (1 - \delta)^3 + R = 1 - 3\delta + 3\delta^2 + O(\delta^3),$$

$$\text{HHI}^2 = (1 - 2\delta + (1 + c)\delta^2)^2 = 1 - 4\delta + (6 + 2c)\delta^2 + O(\delta^3).$$

Substituting into $g(\delta) \equiv \text{HHI} - 4 \sum_i s_i^3 + 3\text{HHI}^2$, the constant terms cancel ($1 - 4 + 3 = 0$), the linear terms give $-2 - (-12) + (-12) = -2$, and the quadratic terms give $(1 + c) - 12 + 3(6 + 2c) = 7(1 + c)$, so

$$g(\delta) = -2\delta + 7(1 + c)\delta^2 + O(\delta^3),$$

which is strictly negative for every sufficiently small $\delta > 0$ (specifically for $0 < \delta < 2/[7(1 + c)]$ to this order), regardless of c — that is, regardless of how the residual share is distributed among the fringe. Since $\mathbb{E}[\Delta\text{HHI}] = \sigma_\varepsilon^2 g(\delta)$, concentration drifts *downward* in expectation whenever one firm holds nearly the whole market, vanishing exactly at $\delta = 0$ (true monopoly) where the index has reached its ceiling of one and can only fall or stay put. Together with the equal-shares check, this pins down the sign pattern Proposition 1 states: positive at the dispersed end (equal shares, any $N \geq 2$), negative at the concentrated end (near monopoly, any fringe structure), and zero exactly at the two boundaries where there is nothing left to sort ($N = 1$) or nowhere further to rise ($\delta = 0$).

A.2 Exactness of the permutation test (Proposition 3 of §5)

The argument is Fisher's randomization principle [42], in the invariance form given by Lehmann and Romano [73] and surveyed by Ernst [38]: a finite group of relabelings that leaves the null model's distribution unchanged makes the observed data point exchangeable with every relabeling of it, and exchangeability alone — with no assumption on the shape of the underlying noise — delivers an exact p -value. What follows verifies that the neutral share process satisfies this invariance and states the resulting formula with every step shown.

Fix one industry-period's data: shares $s = (s_1, \dots, s_N)$ realized as of the start of the period, growth factors $w = (w_1, \dots, w_N)$ realized over the period, and characteristics $z = (z_1, \dots, z_N)$. Under H_0 (Definition 1 applied to the current period), $w_i = 1 + \varepsilon_i$ with $\varepsilon = (\varepsilon_1, \dots, \varepsilon_N)$ iid

across i and independent of z ; s is fixed as of the start of the period and hence independent of the current period's ε as well. Let S_N denote the group of permutations of $\{1, \dots, N\}$, and for $g \in S_N$ write $z_g = (z_{g(1)}, \dots, z_{g(N)})$ for the relabeled characteristics.

Two exchangeability facts. Two consequences of the Definition are all the argument needs.

(F1) ε_i iid $\Rightarrow$ *W 's own coordinates are exchangeable.* Since $\varepsilon_1, \dots, \varepsilon_N$ are iid, $(\varepsilon_1, \dots, \varepsilon_N) \stackrel{d}{=} (\varepsilon_{h(1)}, \dots, \varepsilon_{h(N)})$ for every $h \in S_N$, and hence $W = (1 + \varepsilon_1, \dots, 1 + \varepsilon_N) \stackrel{d}{=} W_h \equiv (w_{h(1)}, \dots, w_{h(N)})$: relabeling which firm gets which growth shock does not change the joint law of the growth-factor vector.

(F2) $\varepsilon \perp z \Rightarrow$ *conditioning on any fixed characteristics vector leaves W 's law unchanged.* Because ε is independent of z (Definition 1), $\mathcal{L}(W \mid Z = \zeta) = \mathcal{L}(W)$ for every fixed vector ζ , in particular for both the observed $\zeta = z_{\text{obs}}$ and for any relabeling $\zeta = z_{\text{obs},g}$ of it: $\mathcal{L}(W \mid Z = z_{\text{obs}}) = \mathcal{L}(W) = \mathcal{L}(W \mid Z = z_{\text{obs},g})$.

Combining (F1) and (F2): conditional on the observed $Z = z_{\text{obs}}$, W is distributed exactly as its unconditional law, and by (F1) that law is itself exchangeable across coordinates, so

$$W_h \mid \{Z = z_{\text{obs}}\} \stackrel{d}{=} W \mid \{Z = z_{\text{obs}}\} \quad \text{for every } h \in S_N. \quad (*)$$

From (*) to a permutation-invariant statistic. Let $T = T(s, w, z)$ be any statistic that depends on the data only through s , w , and the pairing of z with (s, w) — concretely, any statistic built from a sum, over firms i , of a function of (s_i, w_i, z_i) , which is exactly the form of $\text{Sel} = \text{Cov}_s(w, z)$, of $\hat{\beta}$ (an OLS slope, itself a ratio of such sums), and of $T_{\text{Sel}} = \text{Sel}/\hat{\sigma}_{\text{Sel}}$. For such T , relabeling z by g while holding (s, w) fixed produces the same set of firm-level pairs as relabeling w by g^{-1} while holding (s, z) fixed:

$$T(s, w, z_g) = T(s, w_{g^{-1}}, z),$$

a purely notational identity (both sides sum the same N triples (s_i, w_i, z_i) , just indexed differently). Taking $h = g^{-1}$ in (*),

$$T(s, w_{g^{-1}}, z) \mid \{Z = z_{\text{obs}}\} \stackrel{d}{=} T(s, W, z) \mid \{Z = z_{\text{obs}}\} \quad \text{for every } g \in S_N,$$

so that, holding the observed z_{obs} fixed and letting W range over its conditional law, *every one* of the $N!$ recomputed statistics $T(s, W, z_{\text{obs},g})$, $g \in S_N$, has exactly the same distribution: the true null distribution of T .

Exchangeability of the reference set, and the p -value. The observed statistic $T_{\text{obs}} \equiv T(s, w_{\text{obs}}, z_{\text{obs}})$ is the realization, at the single draw $W = w_{\text{obs}}$, of the $g = \text{id}$ member of this family. Draw M permutations $g_1, \dots, g_M$ independently and uniformly from S_N , independently of the data, and set $T_j \equiv T(s, w_{\text{obs}}, z_{\text{obs},g_j})$. Because each T_j is a recomputation of the identity in the family above using only the fixed g_j and the single realized w_{obs} , and because the g_j are chosen without reference to w_{obs} or z_{obs} , the $M + 1$ values $\{T_{\text{obs}}, T_1, \dots, T_M\}$ are exchangeable under H_0 : no member of this set is distinguished from any other by the argument above, which treats $g = \text{id}$ and every sampled g_j identically. For an exchangeable set of $M + 1$ real numbers

(ties broken by a fixed convention, e.g. counting T_{obs} among the values at least as large as itself), the rank of any distinguished member from the top — here, $\text{rank}(T_{\text{obs}}) = 1 + \#\{j : T_j \geq T_{\text{obs}}\}$ — is uniformly distributed on $\{1, \dots, M + 1\}$. Consequently

$$p_{\text{perm}} = \frac{1 + \#\{j : T_j \geq T_{\text{obs}}\}}{1 + M} = \frac{\text{rank}(T_{\text{obs}})}{M + 1}$$

satisfies $\Pr[p_{\text{perm}} \leq \alpha] = \lfloor \alpha(M + 1) \rfloor / (M + 1) \leq \alpha$ for every $\alpha \in (0, 1]$, under H_0 , exactly — not asymptotically, and with no distributional assumption on ε beyond the iid-ness and independence from z that Definition 1 already imposes. Taking $M = N! - 1$ and enumerating every nontrivial permutation recovers the same formula as full enumeration over S_N ; taking $M = 10,000$, as in Proposition 3’s protocol, trades a negligible loss of resolution in α for computational feasibility, with the exactness argument above applying verbatim to whatever M is chosen, since it never used $M = N! - 1$ specifically. $\square$

Why $T_{\mathcal{W}}$ resists this construction. The argument rests entirely on T depending on the data through the pairing of z with (s, w) — formally, on $T(s, w, z_g) \neq T(s, w, z)$ being possible at all for some g , so that recomputation over g generates a nondegenerate reference set. $\mathcal{W} = \sum_i s_i(t+1) \ln w_i$ contains no z_i anywhere in its definition, so $T(s, w, z_g) = T(s, w, z)$ for *every* $g \in S_N$: the “permutation distribution” of $T_{\mathcal{W}}$ is a point mass at the observed value, which carries no information about how unusual that value is. This is not a gap in the argument above but a direct consequence of it, and it is exactly why Proposition 3(iii) requires a bootstrap over ε [31, 37] rather than a relabeling of z : the object whose randomness must be resampled to build $T_{\mathcal{W}}$ ’s null distribution is the noise ε itself, not the pairing of z with anything.

B The biological tests, translated

This appendix states, for readers without a population-genetics background, the five biological test statistics behind Table 1 of Section 2, together with each one’s economic analogue as licensed by this paper’s assembled null (§4: the neutral share process of the Definition, Proposition 1’s concentration drift, and Proposition 2’s restatement of Papers 1, 3 and 7 in the parameters σ_ε and HHI). None of the five biological statistics is new; the translation is.

Ewens’s sampling formula and Watterson’s θ . For a sample of n gene copies drawn from a population evolving under the neutral, infinite-alleles model with a single population-scaled mutation parameter θ , Ewens showed that the probability of any observed partition of the sample into allelic classes — a_j distinct alleles each represented exactly j times, for $j = 1, \dots, n$ — is a known function of n and θ alone [41]. The same θ can be recovered independently from the number of segregating sites S observed in a sample of n sequences, via Watterson’s estimator $\hat{\theta}_W = S/a_n$ with $a_n = \sum_{i=1}^{n-1} 1/i$ [110]. Both routes recover the same drift-and-mutation parameter because, under strict neutrality, that parameter is the only thing generating the data.

Economic analogue. Under the neutral share process, the single parameter governing the entire distribution of outcomes is the growth-shock variance σ_ε^2 together with the number of firms N — exactly the role θ plays for a genetic sample. $N_e = 1/\text{HHI}$ is the market’s Watterson-type estimator: it recovers, from the shares alone, the same “effective number of independent

competitors” that a partition of allele copies recovers for a population. Proposition 1’s expected path for HHI, $\mathbb{E}[\Delta\text{HHI}] = \sigma_\varepsilon^2(\text{HHI} - 4\sum_i s_i^3 + 3\text{HHI}^2)$, is the market’s Ewens-type sampling prediction: the exact distribution concentration should follow given N firms and σ_ε with no fitness differences whatsoever. A market’s N_e falling over time is read correctly only against this sampling prediction, not against the raw number of firms — a N_e that falls exactly as fast as σ_ε^2 and the starting shares predict is evidence of nothing.

Tajima’s D . Tajima’s statistic is built from the relationship between two independent estimators of the population-scaled mutation parameter: Watterson’s $\hat{\theta}_W$, built from the count of segregating sites, and a second estimator built from the average number of pairwise nucleotide differences between sequences, which weights sites by how common the differing base actually is in the sample [108]. Under strict neutrality and constant population size the two estimators agree in expectation, so their standardized difference has expectation zero; population growth or contraction, population structure, or selection all move one estimator relative to the other and displace the statistic from zero. As noted in §2, the verification behind this paper could not confirm Tajima’s own defining formula from the primary 1989 text, so this description is kept at the level the field agrees on rather than presented as a quotation from the source.

Economic analogue. T_W , defined in Proposition 3 as $(W - \frac{\hat{\sigma}_\varepsilon^2}{2}(1 - \text{HHI}))/\hat{\sigma}_W$, is built the same way: it is the standardized difference between the reallocation work actually realized in a market-share panel and the value $\frac{\sigma_\varepsilon^2}{2}(1 - \text{HHI})$ that Proposition 2(ii) says pure noise should produce. Like Tajima’s D , it is a comparison of a realized statistic to what the same neutral parameters should generate, not a comparison of two arbitrarily chosen estimators; and like Tajima’s D , its exact null distribution is not obtained by permutation, because W does not depend on the pairing of z with (s, w) that a relabeling permutes (Proposition 3(iii)). Its null must instead be simulated or bootstrapped over the growth-shock distribution, exactly as Tajima’s D is referred to its null by coalescent simulation rather than by an algebraic permutation argument, using the machinery of Efron and Tibshirani and of Davison and Hinkley [31, 37].

The Hudson–Kreitman–Aguadé (HKA) test. The HKA test generalizes the single-locus comparison to several loci at once. Under strict neutrality with locus-specific mutation rates, the ratio of within-species polymorphism to between-species divergence at a locus should be constant across loci, up to that locus’s own mutation rate; HKA forms a joint statistic testing this constancy over the whole set of loci sampled, and its authors found, applying it to multiple sites in and around the *Adh* region of *Drosophila*, a departure from the constancy the neutral model predicts [61].

Economic analogue. An industry is a locus. Proposition 3’s fifth step — compute p_{perm} industry by industry and combine the results “by a Fisher or Stouffer method” — is the economic HKA: a joint test, across several industries, of whether the ratio of selection-term variance to drift-term variance implied by Sel and σ_ε^2 is the single constant that neutrality with one $(\sigma_\varepsilon, \text{HHI})$ pair per industry predicts, rather than a family of industry-specific departures. Just as HKA needs several loci to be informative, a single industry’s permutation test cannot distinguish a single unlucky draw from a genuine, economy-wide selective pressure; only the cross-industry combination can.

The McDonald–Kreitman test. The McDonald–Kreitman test needs only one locus and no assumption about population size. Substitutions are classified twice over: as replacement (amino-acid-changing) or synonymous, and as a polymorphism segregating within species or a fixed difference between species. Under neutrality the replacement-to-synonymous ratio should be the same whether one looks at polymorphisms or at fixed differences; at the Adh locus, McDonald and Kreitman found instead that “there are more fixed replacement differences between species than expected” [84], reading the excess as adaptive protein evolution rather than drift.

Economic analogue. The 2×2 logic reappears directly in Proposition 3’s permutation test, run on a single industry-period with no assumption about N needed beyond what the permutation itself conditions on. Classify each firm’s growth as linked to its productivity z_i or not, and classify the comparison as the observed pairing of z with (s, w) or as a relabeled one: neutrality (Proposition 3’s H_0) requires the observed pairing to look like a typical relabeling, exactly as it requires the replacement-to-synonymous ratio to look the same across polymorphism and divergence. An observed Sel, or $\hat{\beta}$, in the tail of the permutation distribution is the market’s excess of replacement fixation: a pairing of productivity with growth too tight to be one of the $\binom{N}{\cdot}$ ways the labels could have fallen by chance.

The drift barrier. Ohta’s nearly neutral theory and Lynch’s later formalization agree on a floor: selection cannot reliably act on an effect whose selection coefficient s is smaller than about $1/N_e$, because below that scale genetic drift dominates the outcome regardless of the sign or size of s [80, 92]. The floor rises as the effective population shrinks: a difference that would be decisive in a large population is invisible in a small one.

Economic analogue. Proposition 4’s corollary states the same floor for markets: a productivity difference δz is invisible to market selection when $\beta \delta z$ is small against the drift scale $\sigma_\varepsilon \sqrt{\text{HHI}}$, so that

$$\delta z_{\min} \propto \frac{\sigma_\varepsilon \sqrt{\text{HHI}}}{\beta}.$$

Because $N_e = 1/\text{HHI}$, a concentrated industry is a small- N_e population in exactly Ohta and Lynch’s sense, and it tolerates larger inefficiency differences without selecting on them than a fragmented one does at the same σ_ε and β . Section 6 returns to this floor with the power surface needed to say, for a given industry’s N , HHI and σ_ε , how large δz would have to be before the market could see it at all.

C Simulation: concentration drift, the permutation null, and power

A deterministic simulation (script `sim_null.py`, seed 20260830) checks Proposition 1, stages the permutation test of Proposition 3, and computes the power surface of Proposition 4. Throughout, growth shocks ε_i are Laplace-distributed — the tent-shaped growth-rate distribution of §3 — and the neutral share process is $s_i(t+1) \propto s_i(t)(1 + \varepsilon_i)$.

Concentration drifts upward (Proposition 1). Two industries of $N = 100$ firms are run under pure noise with $\sigma_\varepsilon = 0.10$: one from equal shares ($\text{HHI}_0 = 0.0100$) and one from a Zipf-like rank-size start ($\text{HHI}_0 = 0.0188$). The one-period expectation $\mathbb{E}[\Delta \text{HHI}]$, estimated from

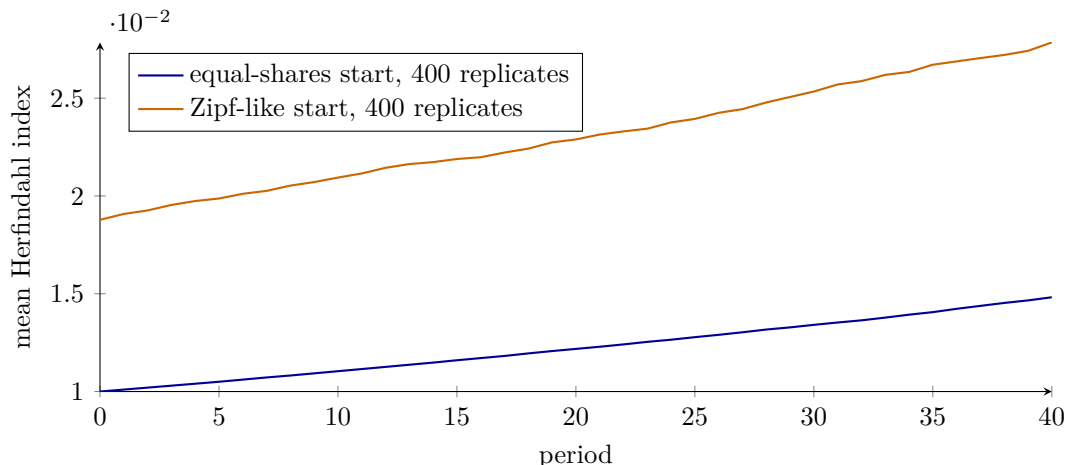


Figure 1: Mean Herfindahl index under the neutral share process ($\sigma_\varepsilon = 0.10$, 400 replicates) from an equal-shares start and from a Zipf-like start. Concentration rises under pure noise at the rate of Proposition 1.

40,000 independent single-step draws, is 0.000099 ± 0.0000004 for the equal start against the formula's 0.000099 (ratio 1.002), and 0.000175 ± 0.000004 for the Zipf start against 0.000170 (ratio 1.029, the residual being the fourth-order term that the Laplace kurtosis inflates). Over forty periods the mean Herfindahl index across 400 replicates rises from 0.0100 to 0.0148 and from 0.0188 to 0.0279 (Figure 1): in forty periods of pure noise the equal-shares industry acquires half again its concentration and the Zipf industry half again its own. No firm is better than any other in either economy.

The permutation null (Proposition 3). One industry of $N = 200$ firms with Zipf-like shares carries log productivities $z_i \sim \mathcal{N}(0, 0.30^2)$ and growth shocks with $\sigma_\varepsilon = 0.20$. Under neutrality ($\beta = 0$) the observed Price selection term is $\text{Sel} = -0.0046$; its permutation distribution over 4,000 reshufflings of the productivity labels has standard deviation 0.0050, so $T_{\text{Sel}} = -0.91$ and $p_{\text{perm}} = 0.82$ — not rejected, correctly. With selection at $\beta = 0.10$ added to the same industry the observed term is $+0.0165$, $T_{\text{Sel}} = 1.26$ and $p_{\text{perm}} = 0.10$ — *not* rejected either: run once on one industry of two hundred firms, the test cannot see a ten-percent selection intensity through growth noise of this size, which is what the power surface below says it should not be expected to do (at $N = 200$ the power against $\beta = 0.10$ is 0.40 with $\sigma_\varepsilon = 0.30$ and 0.98 with $\sigma_\varepsilon = 0.10$; $\sigma_\varepsilon = 0.20$ sits between). With selection at $\beta = 0.20$ — a separate draw of the shocks on its own generator, so that the rest of the appendix is unaffected — the observed term is $+0.0161$, $T_{\text{Sel}} = 2.94$ and $p_{\text{perm}} = 0.002$ — rejected (Figure 2). The permutation distribution is the exact conditional null: it holds the shares, the growth shocks and the productivities fixed and randomizes only their pairing, which is the only thing neutrality restricts. Note what the two selective panels show together: an observed selection term of 0.0165 was not evidence of selection, and an observed term of 0.0161 was, because the null distributions behind them differed by a factor of two and a half in width. The observed number alone carries no verdict; the verdict is a property of the number and its null.

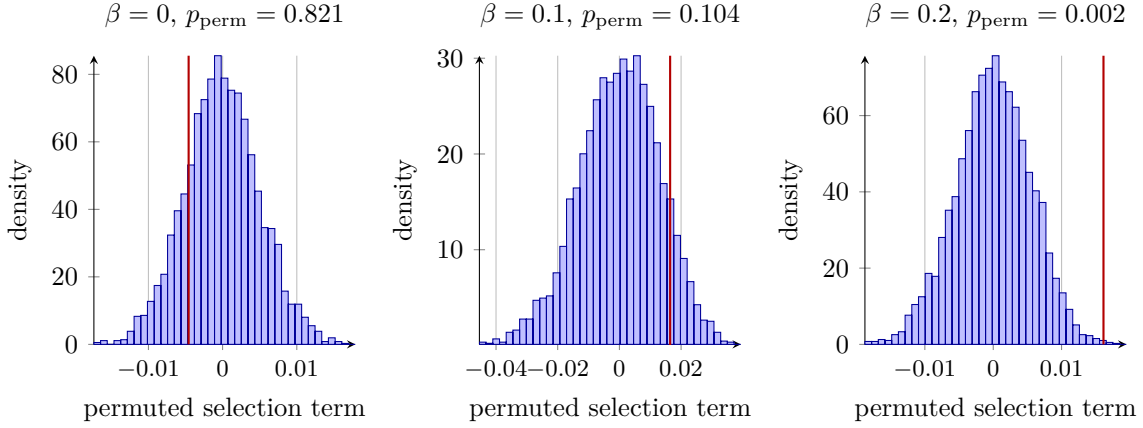


Figure 2: Permutation null distributions of the Price selection term for one industry of 200 firms without selection (left), with $\beta = 0.10$ (centre) and with $\beta = 0.20$ (right); the observed value is the vertical line. The neutral draw is inside its null; the $\beta = 0.10$ draw is positive but also inside its null; only the $\beta = 0.20$ draw is in the tail.

Power (Proposition 4). For each combination of $N \in \{50, 200, 1000\}$, $\sigma_\varepsilon \in \{0.10, 0.30\}$, $\beta \in \{0, 0.02, 0.05, 0.10, 0.20\}$ and share structure (equal, or Zipf-like with $\text{HHI} = 0.033, 0.011, 0.003$ at the three sizes), 200 industries are simulated with $\sigma_z = 0.30$ and the permutation test at size 0.05 is run with 199 permutations; power is the rejection frequency (Table 5, Figure 3). The test holds its size: at $\beta = 0$ the rejection rate lies between 0.04 and 0.08 in every cell. Against selection, the findings are the ones Proposition 4 asserts. With mild growth noise ($\sigma_\varepsilon = 0.10$) the test is powerful: at $N = 200$ it detects $\beta = 0.05$ about two-thirds of the time and $\beta = 0.10$ almost always. With growth noise at the Laplace scale of real firm panels ($\sigma_\varepsilon = 0.30$) it is not: at $N = 200$ the power against $\beta = 0.10$ is 0.40 with equal shares and 0.28 with Zipf-like shares, and against $\beta = 0.05$ it is 0.14 under either — a test that would miss a selection intensity of five percent per period six times out of seven. Only at $N = 1000$ does power against $\beta = 0.10$ reach 0.93 (equal shares) — and 0.56 with Zipf-like shares, because concentration reduces the effective number of independent comparisons exactly as the drift scale $\sigma_\varepsilon \sqrt{\text{HHI}}$ of Proposition 4’s corollary says.

The re-reading of §7 rests on this table. Where a published study reports a small, marginally significant productivity–growth coefficient from an industry of a few hundred firms with growth dispersion at the Laplace scale, the table says that the study was operating at a power of between one in seven and two in five against the effect sizes it reports; its finding is consistent with weak selection, and equally consistent with drift and a lucky draw, and the two cannot be told apart without the null.

References

- [1] David Alonso, Rampal S. Etienne, Alan J. McKane (2006). The merits of neutral theory. *Trends in Ecology and Evolution*, vol. 21, no. 8, pp. 451–457. <https://pubmed.ncbi.nlm.nih.gov/16766082/>. DOI: 10.1016/j.j.tree.2006.03.019.
- [2] Luís A. Nunes Amaral, Sergey V. Buldyrev, Shlomo Havlin, Heiko Leschhorn, Philipp

Shares	σ_ε	N	$\beta = 0$	0.02	0.05	0.10	0.20
equal	0.10	50	0.06	0.06	0.23	0.60	0.98
equal	0.10	200	0.06	0.23	0.65	0.98	1.00
equal	0.10	1000	0.04	0.58	1.00	1.00	1.00
equal	0.30	50	0.05	0.05	0.12	0.15	0.40
equal	0.30	200	0.04	0.07	0.14	0.40	0.86
equal	0.30	1000	0.05	0.10	0.47	0.93	1.00
Zipf-like	0.10	50	0.08	0.09	0.20	0.52	0.88
Zipf-like	0.10	200	0.07	0.14	0.42	0.91	1.00
Zipf-like	0.10	1000	0.06	0.28	0.84	1.00	1.00
Zipf-like	0.30	50	0.04	0.06	0.04	0.12	0.37
Zipf-like	0.30	200	0.04	0.06	0.14	0.28	0.65
Zipf-like	0.30	1000	0.04	0.07	0.28	0.56	0.98

Table 5: Power of the permutation T_{Sel} test at size 0.05 (200 simulated industries per cell, 199 permutations each, $\sigma_z = 0.30$, Laplace growth shocks). The first column of results is the realized size.

- Maass, Michael A. Salinger, H. Eugene Stanley, Michael H. R. Stanley (1997). Scaling Behavior in Economics: I. Empirical Results for Company Growth. *Journal de Physique I (France)*, vol. 7. <https://arxiv.org/abs/cond-mat/9702082>.
- [3] Esben Sloth Andersen (2004). Population Thinking, Price’s Equation and the Analysis of Economic Evolution. *Evolutionary and Institutional Economics Review*, Vol. 1, No. 1, pp. 127-148. <https://vbn.aau.dk/en/publications/population-thinking-prices-equation-and-the-analysis-of-economic-/>. DOI: 10.14441/eier.1.127.
- [4] Yoshiyuki Arata (2019). Firm Growth and Laplace Distribution: The Importance of Large Jumps. *Journal of Economic Dynamics and Control*, vol. 103, pp. 63-82. <https://www.rieti.go.jp/jp/publications/dp/14e033.pdf>. DOI: 10.1016/j.jedc.2019.01.009.
- [5] Robert L. Axtell (2001). Zipf Distribution of U.S. Firm Sizes. *Science*, vol. 293, no. 5536. <https://faculty.sites.iastate.edu/tesfatsi/archive/tesfatsi/ZipfDistributionFirmSizes.RAxtell2001.pdf>. DOI: 10.1126/science.1062081.
- [6] Martin Neil Baily, Charles Hulten, David Campbell (1992). Productivity Dynamics in Manufacturing Plants. *Brookings Papers on Economic Activity: Microeconomics*, 1992. <https://www.brookings.edu/articles/productivity-dynamics-in-manufacturing-plants/>.
- [7] Eric Bartelsman, John Haltiwanger, Stefano Scarpetta (2013). Cross-Country Differences in Productivity: The Role of Allocation and Selection. *American Economic Review*, Vol. 103, No. 1, pp. 305-334. https://econweb.umd.edu/~haltiwan/Cross_Country_Differences_Productivity_AER_2013.pdf. DOI: 10.1257/aer.103.1.305.
- [8] Graham Bell (2001). Neutral Macroecology. *Science*, vol. 293, no. 5539, pp. 2413-2418. http://www.sfu.ca/biology/courses/bisc830/Bell_2001_Science2.pdf. DOI: 10.1126/science.293.5539.2413.

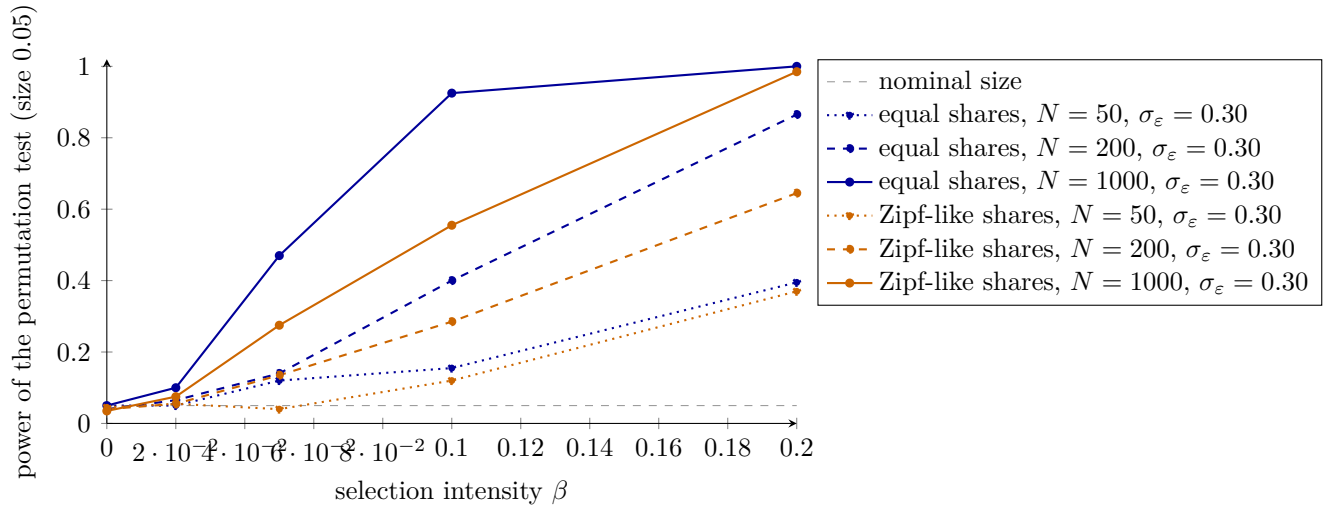


Figure 3: Power of the permutation test against the selection intensity β at $\sigma_\varepsilon = 0.30$, for three industry sizes and two share structures. With realistic growth noise, an industry of a few hundred firms cannot reliably distinguish a five- or ten-percent selection intensity from drift; concentration lowers power further.

- [9] Giulio Bottazzi, Angelo Secchi (2006). Explaining the Distribution of Firm Growth Rates. *The RAND Journal of Economics*, vol. 37, no. 2. <https://www.lem.santannapisa.it/WPLem/files/2005-16.pdf>. DOI: 10.1111/j.1756-2171.2006.tb00014.x.
- [10] Giulio Bottazzi, Giovanni Dosi, Nadia Jacoby, Angelo Secchi, Federico Tamagni (2010). Corporate performances and market selection. Some comparative evidence. *Industrial and Corporate Change*, Vol. 19, No. 6, pp. 1953-1996 (opened as LEM Working Paper 2009/13). <https://www.lem.santannapisa.it/WPLem/files/2009-13.pdf>. DOI: 10.1093/icc/dtq063.
- [11] Kenneth P. Burnham, David R. Anderson (2002). Model Selection and Multi-model Inference: A Practical Information-Theoretic Approach (2nd ed.). *Springer-Verlag, New York/Heidelberg*. <https://link.springer.com/book/10.1007/b97636>. DOI: 10.1007/b97636.
- [12] Luís M. B. Cabral, José Mata (2003). On the Evolution of the Firm Size Distribution: Facts and Theory. *American Economic Review*, vol. 93, no. 4, pp. 1075–1090. <https://ideas.repec.org/a/aea/aecrev/v93y2003i4p1075-1090.html>. DOI: 10.1257/000282803769206205.
- [13] Uwe Cantner, Jens J. Krüger (2008). Micro-heterogeneity and aggregate productivity development in the German manufacturing sector. *Journal of Evolutionary Economics*, Vol. 18, No. 2, pp. 119-133. <https://link.springer.com/article/10.1007/s00191-007-0079-z>. DOI: 10.1007/s00191-007-0079-z.
- [14] D. G. Champernowne (1953). A Model of Income Distribution. *The Economic Journal*,

- vol. 63, no. 250. <https://academic.oup.com/ej/article/63/250/318/5258723>. DOI: 10.2307/2227127.
- [15] Brian Charlesworth (2009). Effective Population Size and Patterns of Molecular Evolution and Variation. *Nature Reviews Genetics*, vol. 10, pp. 195–205. <https://www.nature.com/articles/nrg2526>. DOI: 10.1038/nrg2526.
 - [16] Jérôme Chave (2004). Neutral theory and community ecology. *Ecology Letters*, vol. 7, no. 3, pp. 241–253. <https://kevintshoemaker.github.io/EECB-703/Chave%20-%202004%20-%20Neutral%20theory%20and%20community%20ecology%20Neutral%20theo.pdf>. DOI: 10.1111/j.1461-0248.2003.00566.x.
 - [17] Aaron Clauset, Cosma Rohilla Shalizi, M. E. J. Newman (2009). Power-Law Distributions in Empirical Data. *SIAM Review*, vol. 51, no. 4, pp. 661–703. <https://arxiv.org/abs/0706.1062>. DOI: 10.1137/070710111.
 - [18] Alex Coad (2007). Testing the principle of “growth of the fitter”: The relationship between profits and firm growth. *Structural Change and Economic Dynamics*, Vol. 18, No. 3, pp. 370–386. <https://ideas.repec.org/a/eee/streco/v18y2007i3p370-386.html>. DOI: 10.1016/j.strueco.2007.05.001.
 - [19] Alex Coad (2009). The Growth of Firms: A Survey of Theories and Empirical Evidence. *Edward Elgar Publishing (Cheltenham, UK)*. <https://archive.org/details/growthoffirmssur0000coad>.
 - [20] Alex Coad, Agustí Segarra, Mercedes Teruel (2013). Like milk or wine: Does firm performance improve with age? *Structural Change and Economic Dynamics*, Vol. 24, pp. 173–189 (opened as XREAP Working Paper 2010-10, “Document de Treball XREAP2010-10”). <http://www.xreap.cat/RePEc/xrp/pdf/XREAP2010-10.pdf>. DOI: 10.1016/j.strueco.2012.07.002.
 - [21] Rama Cont, Didier Sornette (1997). Convergent Multiplicative Processes Repelled from Zero: Power Laws and Truncated Power Laws. *Journal de Physique I*, vol. 7, no. 3, pp. 431–444. <https://arxiv.org/abs/cond-mat/9609074>. DOI: 10.1051/jp1:1997169.
 - [22] Brecht Corbeel (2026). The Cooperative Firm: Multilevel Selection Inside Organizations. *The Evolutionary Economics Papers*, No. 5, Absolute Digital Publishers. <https://absolutedigitalpublishers.com/papers/evolutionary-economics-05-cooperative-firm.pdf>.
 - [23] Brecht Corbeel (2026). Frozen Accidents: The Population Genetics of Standards. *The Evolutionary Economics Papers*, No. 7, Absolute Digital Publishers. <https://absolutedigitalpublishers.com/papers/evolutionary-economics-07-frozen-accidents.pdf>.
 - [24] Brecht Corbeel (2026). Money: The Convention That Selected Itself. *The Evolutionary Economics Papers*, No. 8, Absolute Digital Publishers. <https://absolutedigitalpublishers.com/papers/evolutionary-economics-08-money-selected-itself.pdf>.

- [25] Brecht Corbeel (2026). The Red Queen Economy: Why Profit Refuses to Persist. *The Evolutionary Economics Papers*, No. 9, Absolute Digital Publishers. <https://absolutedigitalpublishers.com/papers/evolutionary-economics-09-red-queen-economy.pdf>.
- [26] Brecht Corbeel (2026). Replicator Markets: When Competition Computes. *The Evolutionary Economics Papers*, No. 2, Absolute Digital Publishers. <https://absolutedigitalpublishers.com/papers/evolutionary-economics-02-replicator-markets.pdf>.
- [27] Brecht Corbeel (2026). Routines Are Genes: Nelson and Winter Formalized Forty Years Late. *The Evolutionary Economics Papers*, No. 6, Absolute Digital Publishers. <https://absolutedigitalpublishers.com/papers/evolutionary-economics-06-routines-are-genes.pdf>.
- [28] Brecht Corbeel (2026). Rugged Money: NK Landscapes and the Search for Better Technology. *The Evolutionary Economics Papers*, No. 4, Absolute Digital Publishers. <https://absolutedigitalpublishers.com/papers/evolutionary-economics-04-rugged-landscapes.pdf>.
- [29] Brecht Corbeel (2026). Schumpeter's Demon: Entry, Exit, and the Thermodynamics of Creative Destruction. *The Evolutionary Economics Papers*, No. 3, Absolute Digital Publishers. <https://absolutedigitalpublishers.com/papers/evolutionary-economics-03-schumpeters-demon.pdf>.
- [30] Brecht Corbeel (2026). Selection Accounting: The Price Equation as the Fundamental Theorem of Evolutionary Economics. *The Evolutionary Economics Papers*, No. 1, Absolute Digital Publishers. <https://absolutedigitalpublishers.com/papers/evolutionary-economics-01-selection-accounting.pdf>.
- [31] A.C. Davison, D.V. Hinkley (1997). Bootstrap Methods and Their Application. *Cambridge University Press (Cambridge Series in Statistical and Probabilistic Mathematics)*. <https://doi.org/10.1017/cbo9780511802843>. DOI: 10.1017/cbo9780511802843.
- [32] Ryan A. Decker, John Haltiwanger, Ron S. Jarmin, Javier Miranda (2020). Changing Business Dynamism and Productivity: Shocks versus Responsiveness. *American Economic Review*, Vol. 110, No. 12, pp. 3952-3990 (circulated as NBER Working Paper 24236, Jan. 2018). <https://www.nber.org/papers/w24236>. DOI: 10.1257/aer.20190680.
- [33] Giovanni Dosi, Daniele Moschella, Emanuele Pugliese, Federico Tamagni (2015). Productivity, market selection and corporate growth: comparative evidence across US and Europe. *Small Business Economics*, Vol. 45, No. 3, pp. 643-672 (opened as LEM Working Paper 2013/15, "Forthcoming in Small Business Economics"). <https://www.lem.santannapisa.it/WPLem/files/2013-15.pdf>. DOI: 10.1007/s11187-015-9655-z.
- [34] Giovanni Dosi, Marco Grazzi, Daniele Moschella, Gary Pisano, Federico Tamagni (2020). Long-Term Firm Growth: An Empirical Analysis of US Manufacturers 1959-2015. *Industrial and Corporate Change*, Vol. 29, No. 2, pp. 309-332 (opened as LEM Working Paper 2019/13). <https://www.lem.santannapisa.it/WPLem/files/2019-13.pdf>. DOI: 10.1093/icc/dtz044.

- [35] Timothy Dunne, Mark J. Roberts, Larry Samuelson (1989). The Growth and Failure of U.S. Manufacturing Plants. *The Quarterly Journal of Economics*, vol. 104, no. 4. <https://doi.org/10.2307/2937862>. DOI: 10.2307/2937862.
- [36] Jan Eeckhout (2004). Gibrat's Law for (All) Cities. *American Economic Review*, vol. 94, no. 5, pp. 1429–1451. <https://ideas.repec.org/a/aea/aecrev/v94y2004i5p1429-1451.html>. DOI: 10.1257/0002828043052303.
- [37] Bradley Efron, Robert J. Tibshirani (1993). An Introduction to the Bootstrap. *Chapman & Hall, New York (Monographs on Statistics and Applied Probability)*. <https://archive.org/details/introductiontobo0000efro>. DOI: 10.1007/978-1-4899-4541-9.
- [38] Michael D. Ernst (2004). Permutation Methods: A Basis for Exact Inference. *Statistical Science*, vol. 19, no. 4, pp. 676–685. <https://projecteuclid.org/journals/statistical-science/volume-19/issue-4/Permutation-Methods-A-Basis-for-Exact-Inference/10.1214/088342304000000396.full>. DOI: 10.1214/088342304000000396.
- [39] Rampal S. Etienne (2005). A new sampling formula for neutral biodiversity. *Ecology Letters*, vol. 8, no. 3, pp. 253–260. <https://onlinelibrary.wiley.com/doi/abs/10.1111/j.1461-0248.2004.00717.x>. DOI: 10.1111/j.1461-0248.2004.00717.x.
- [40] David S. Evans (1987). Tests of Alternative Theories of Firm Growth. *Journal of Political Economy*, vol. 95, no. 4. <https://doi.org/10.1086/261480>. DOI: 10.1086/261480.
- [41] Warren J. Ewens (1972). The Sampling Theory of Selectively Neutral Alleles. *Theoretical Population Biology*, vol. 3, pp. 87–112. [https://doi.org/10.1016/0040-5809\(72\)90035-4](https://doi.org/10.1016/0040-5809(72)90035-4). DOI: 10.1016/0040-5809(72)90035-4.
- [42] Ronald A. Fisher (1935). The Design of Experiments. *Oliver and Boyd, Edinburgh/London (book)*. <https://archive.org/details/in.ernet.dli.2015.502684>.
- [43] Lucia Foster, John Haltiwanger, C.J. Krizan (2001). Aggregate Productivity Growth: Lessons from Microeconomic Evidence. *Chapter 8 in “New Developments in Productivity Analysis” (C.R. Hulten, E.R. Dean, M.J. Harper, eds.)*, University of Chicago Press for NBER, pp. 303–372 (circulated as NBER Working Paper 6803). <https://www.nber.org/system/files/chapters/c10129/c10129.pdf>.
- [44] Lucia Foster, John Haltiwanger, Chad Syverson (2008). Reallocation, Firm Turnover, and Efficiency: Selection on Productivity or Profitability? *American Economic Review*, Vol. 98, No. 1, pp. 394–425 (opened as NBER Working Paper 11555, Aug. 2005). <https://www.nber.org/papers/w11555>. DOI: 10.1257/aer.98.1.394.
- [45] Steven A. Frank (2012). Natural selection. IV. The Price equation. *Journal of Evolutionary Biology*, Vol. 25, No. 6, pp. 1002–1019. <https://arxiv.org/abs/1204.1515>. DOI: 10.1111/j.1420-9101.2012.02498.x.
- [46] Dongfeng Fu, Fabio Pammolli, Sergey V. Buldyrev, Massimo Riccaboni, Kaushik Matia, Kazuko Yamasaki, H. Eugene Stanley (2005). The Growth of Business Firms: Theoretical Framework and Empirical Evidence. *Proceedings of the National Academy of Sciences*

- (PNAS), vol. 102, no. 52. <https://pmc.ncbi.nlm.nih.gov/articles/PMC1323202/>. DOI: 10.1073/pnas.0509543102.
- [47] Xavier Gabaix (1999). Zipf’s Law for Cities: An Explanation. *Quarterly Journal of Economics*, vol. 114, pp. 739–767. https://econpapers.repec.org/article/oupqjecon/v_3a114_3ay_3a1999_3ai_3a3_3ap_3a739-767..htm. DOI: 10.1162/003355399556133.
 - [48] Xavier Gabaix (2009). Power Laws in Economics and Finance. *Annual Review of Economics*, vol. 1, pp. 255–294 (NBER Working Paper 14299 open-access predecessor). <https://www.nber.org/papers/w14299>. DOI: 10.1146/annurev.economics.050708.142940.
 - [49] Robert Gibrat (1931). Les inégalités économiques: applications aux inégalités des richesses, à la concentration des entreprises, aux populations des villes, aux statistiques des familles, etc., d’une loi nouvelle, la loi de l’effet proportionnel. *Librairie du Recueil Sirey (Paris) — doctoral thesis, University of Lyon*. <https://search.worldcat.org/title/inegalites-economiques/oclc/250098599>.
 - [50] John H. Gillespie (1991). The Causes of Molecular Evolution. *Oxford University Press, Oxford Series in Ecology and Evolution (book)*. <https://openlibrary.org/works/OL2673830W>.
 - [51] Phillip I. Good (2005). Permutation, Parametric, and Bootstrap Tests of Hypotheses (3rd ed.). *Springer (Springer Series in Statistics), New York*. <https://doi.org/10.1007/b138696>. DOI: 10.1007/b138696.
 - [52] Nicholas J. Gotelli, Gary R. Graves (1996). Null Models in Ecology. *Smithsonian Institution Press (book), Washington*. <https://archive.org/details/nullmodelsinecol0000gote>.
 - [53] Matthew W. Hahn (2008). Toward a Selection Theory of Molecular Evolution. *Evolution*, vol. 62, no. 2, pp. 255–265. <https://academic.oup.com/evolut/article/62/2/255/6854555>. DOI: 10.1111/j.1558-5646.2007.00308.x.
 - [54] Bronwyn H. Hall (1987). The Relationship Between Firm Size and Firm Growth in the U.S. Manufacturing Sector. *The Journal of Industrial Economics*, vol. 35, no. 4 (also circulated as NBER Working Paper No. 1965). <https://doi.org/10.2307/2098589>. DOI: 10.2307/2098589.
 - [55] P. E. Hart, S. J. Prais (1956). The Analysis of Business Concentration: A Statistical Approach. *Journal of the Royal Statistical Society, Series A (General)*, vol. 119, no. 2. <https://doi.org/10.2307/2342882>. DOI: 10.2307/2342882.
 - [56] Paul H. Harvey, Robert K. Colwell, Jonathan W. Silvertown, Robert M. May (1983). Null Models in Ecology. *Annual Review of Ecology and Systematics*, vol. 14, pp. 189–211. <https://doi.org/10.1146/annurev.es.14.110183.001201>. DOI: 10.1146/annurev.es.14.110183.001201.
 - [57] Chang-Tai Hsieh, Peter J. Klenow (2009). Misallocation and Manufacturing TFP in China and India. *Quarterly Journal of Economics*, Vol. 124, No. 4, pp. 1403–1448 (circulated as NBER Working Paper 13290, Aug. 2007, revised Dec. 2008). <https://www.nber.org/papers/w13290>. DOI: 10.1162/qjec.2009.124.4.1403.

- [58] Stephen P. Hubbell (1979). Tree Dispersion, Abundance, and Diversity in a Tropical Dry Forest. *Science*, vol. 203, no. 4387, pp. 1299–1309. <https://pubmed.ncbi.nlm.nih.gov/17780463/>. DOI: 10.1126/science.203.4387.1299.
- [59] Stephen P. Hubbell (2001). The Unified Neutral Theory of Biodiversity and Biogeography. *Princeton University Press (book; Monographs in Population Biology series)*. <https://press.princeton.edu/books/paperback/9780691021287/the-unified-neutral-theory-of-biodiversity-and-biogeography>.
- [60] Stephen P. Hubbell (2005). Neutral theory in community ecology and the hypothesis of functional equivalence. *Functional Ecology*, vol. 19, pp. 166–172. <https://people.uncw.edu/borretts/courses/BI0602/Hubbell1%202005%20neutral%20theory%20in%20community%20ecology%20and%20the%20hypothesis%20of%20functional%20equivalence.pdf>. DOI: 10.1111/j.0269-8463.2005.00965.x.
- [61] Richard R. Hudson, Martin Kreitman, Montserrat Aguadé (1987). A Test of Neutral Molecular Evolution Based on Nucleotide Data. *Genetics*, vol. 116, no. 1, pp. 153–159. <https://pmc.ncbi.nlm.nih.gov/articles/PMC1203113/>. DOI: 10.1093/genetics/116.1.153.
- [62] Yuji Ijiri, Herbert A. Simon (1977). Skew Distributions and the Sizes of Business Firms (Contributions to Economic Analysis, vol. 24). *North-Holland Publishing Company (Amsterdam/New York)*. https://books.google.com/books/about/Skew_Distributions_and_the_Sizes_of_Busi.html?id=1_K3AAAAIAAJ.
- [63] Jeffrey D. Jensen, Bret A. Payseur, Wolfgang Stephan, Charles F. Aquadro, Michael Lynch, Deborah Charlesworth, Brian Charlesworth (2019). The Importance of the Neutral Theory in 1968 and 50 Years On: A Response to Kern and Hahn 2018. *Evolution*, vol. 73, no. 1, pp. 111–114. <https://pmc.ncbi.nlm.nih.gov/articles/PMC6496948/>. DOI: 10.1111/evo.13650.
- [64] M. Kalecki (1945). On the Gibrat Distribution. *Econometrica*, vol. 13, no. 2. <https://doi.org/10.2307/1907013>. DOI: 10.2307/1907013.
- [65] Andrew D. Kern, Matthew W. Hahn (2018). The Neutral Theory in Light of Natural Selection. *Molecular Biology and Evolution*, vol. 35, no. 6, pp. 1366–1371. <https://pmc.ncbi.nlm.nih.gov/articles/PMC5967545/>. DOI: 10.1093/molbev/msy092.
- [66] Harry Kesten (1973). Random difference equations and renewal theory for products of random matrices. *Acta Mathematica*, vol. 131, pp. 207–248. <https://projecteuclid.org/euclid.acta/1485889791>. DOI: 10.1007/BF02392040.
- [67] Motoo Kimura (1962). On the Probability of Fixation of Mutant Genes in a Population. *Genetics*, vol. 47, no. 6, pp. 713–719. <https://pmc.ncbi.nlm.nih.gov/articles/PMC1210364/>. DOI: 10.1093/genetics/47.6.713.
- [68] Motoo Kimura (1968). Evolutionary Rate at the Molecular Level. *Nature*, vol. 217, pp. 624–626. <https://www.nature.com/articles/217624a0>. DOI: 10.1038/217624a0.
- [69] Motoo Kimura (1983). The Neutral Theory of Molecular Evolution. *Cambridge University Press (book)*. <https://openlibrary.org/works/OL3478974W>.

- [70] Jack L. King, Thomas H. Jukes (1969). Non-Darwinian Evolution. *Science*, vol. 164, no. 3881, pp. 788–798. <https://doi.org/10.1126/science.164.3881.788>. DOI: 10.1126/science.164.3881.788.
- [71] Thorbjørn Knudsen (2002). Economic selection theory. *Journal of Evolutionary Economics*, Vol. 12, No. 4, pp. 443–470. <https://ideas.repec.org/a/spr/joevec/v12y2002i4p443-470.html>. DOI: 10.1007/s00191-002-0126-8.
- [72] Martin Kreitman (2000). Methods to Detect Selection in Populations with Applications to the Human. *Annual Review of Genomics and Human Genetics*, vol. 1, pp. 539–559. <https://doi.org/10.1146/annurev.genom.1.1.539>. DOI: 10.1146/annurev.genom.1.1.539.
- [73] Erich L. Lehmann, Joseph P. Romano (2005). Testing Statistical Hypotheses (3rd ed.). *Springer (Springer Texts in Statistics)*, New York. <https://doi.org/10.1007/0-387-27605-x>. DOI: 10.1007/0-387-27605-x.
- [74] Egbert G. Leigh Jr. (2007). Neutral theory: a historical perspective. *Journal of Evolutionary Biology*, vol. 20, no. 6, pp. 2075–2091. <https://pubmed.ncbi.nlm.nih.gov/17956380/>. DOI: 10.1111/j.1420-9101.2007.01410.x.
- [75] Moshe Levy, Sorin Solomon (1996). Power Laws are Logarithmic Boltzmann Laws. *International Journal of Modern Physics C*, vol. 7, no. 4, pp. 595–601. <https://arxiv.org/abs/adap-org/9607001>. DOI: 10.1142/S0129183196000491.
- [76] Moshe Levy (2009). Gibrat’s Law for (All) Cities: Comment. *American Economic Review*, vol. 99, no. 4, pp. 1672–1675. <https://ideas.repec.org/a/aea/aecrev/v99y2009i4p1672-75.html>. DOI: 10.1257/aer.99.4.1672. And Jan Eeckhout (2009). Gibrat’s Law for (All) Cities: Reply. *American Economic Review*, vol. 99, no. 4, pp. 1676–1683. <https://ideas.repec.org/a/aea/aecrev/v99y2009i4p1676-83.html>. DOI: 10.1257/aer.99.4.1676.
- [77] Francesca Lotti, Enrico Santarelli, Marco Vivarelli (2003). Does Gibrat’s Law Hold Among Young, Small Firms? *Journal of Evolutionary Economics*, vol. 13, no. 3. <https://doi.org/10.1007/s00191-003-0153-0>. DOI: 10.1007/s00191-003-0153-0.
- [78] Erzo G. J. Luttmer (2007). Selection, Growth, and the Size Distribution of Firms. *Quarterly Journal of Economics*, vol. 122, pp. 1103–1144. <https://users.econ.umn.edu/~holmes/class/2007f8601/papers/luttmer.pdf>. DOI: 10.1162/qjec.122.3.1103.
- [79] Michael Lynch, John S. Conery (2003). The Origins of Genome Complexity. *Science*, vol. 302, pp. 1401–1404. <https://doi.org/10.1126/science.1089370>. DOI: 10.1126/science.1089370.
- [80] Michael Lynch (2007). The Frailty of Adaptive Hypotheses for the Origins of Organismal Complexity. *Proceedings of the National Academy of Sciences (PNAS)*, vol. 104, Suppl. 1, pp. 8597–8604. <https://pmc.ncbi.nlm.nih.gov/articles/PMC1876435/>. DOI: 10.1073/pnas.0702207104.
- [81] Michael Lynch (2007). The Origins of Genome Architecture. *Sinauer Associates (book)*. <https://openlibrary.org/works/OL2668954W>.

- [82] Yannick Malevergne, Alexander Saichev, Didier Sornette (2013). Zipf's Law and Maximum Sustainable Growth. *Journal of Economic Dynamics and Control*, vol. 37, no. 6, pp. 1195–1212. <https://arxiv.org/abs/1012.0199>. DOI: 10.1016/j.jedc.2013.02.004.
- [83] Edwin Mansfield (1962). Entry, Gibrat's Law, Innovation, and the Growth of Firms. *American Economic Review*, vol. 52, no. 5 (a 1961 working-paper version circulated as Cowles Foundation Discussion Paper No. 126, Yale). <https://elischolar.library.yale.edu/cowles-discussion-paper-series/355/>.
- [84] John H. McDonald, Martin Kreitman (1991). Adaptive Protein Evolution at the Adh Locus in Drosophila. *Nature*, vol. 351, pp. 652–654. <https://www.nature.com/articles/351652a0>. DOI: 10.1038/351652a0.
- [85] Brian J. McGill (2003). A test of the unified neutral theory of biodiversity. *Nature*, vol. 422, pp. 881–885. <http://brianmcgill.org/zsm.pdf>. DOI: 10.1038/nature01583.
- [86] Marc J. Melitz, Sašo Polanec (2015). Dynamic Olley-Pakes Productivity Decomposition with Entry and Exit. *RAND Journal of Economics*, Vol. 46, No. 2, pp. 362–375 (circulated as NBER Working Paper 18182, issued June 2012, revised September 2014). <https://www.nber.org/papers/w18182>. DOI: 10.1111/1756-2171.12088.
- [87] J. Stanley Metcalfe (1994). Competition, Fisher's Principle and Increasing Returns in the Selection Process. *Journal of Evolutionary Economics*, Vol. 4, No. 4, pp. 327–346. <https://ideas.repec.org/a/spr/joevec/v4y1994i4p327-46.html>. DOI: 10.1007/BF01236409.
- [88] Michael Mitzenmacher (2004). A Brief History of Generative Models for Power Law and Lognormal Distributions. *Internet Mathematics*, vol. 1, no. 2, pp. 226–251. <https://www.internetmathematicsjournal.com/article/1385.pdf>. DOI: 10.1080/15427951.2004.10129088.
- [89] Masatoshi Nei (2005). Selectionism and Neutralism in Molecular Evolution. *Molecular Biology and Evolution*, vol. 22, no. 12, pp. 2318–2342. <https://pmc.ncbi.nlm.nih.gov/articles/PMC1513187/>. DOI: 10.1093/molbev/msi242.
- [90] M. E. J. Newman (2005). Power laws, Pareto distributions and Zipf's law. *Contemporary Physics*, vol. 46, pp. 323–351. <https://arxiv.org/abs/cond-mat/0412004>. DOI: 10.1080/00107510500052444.
- [91] Rasmus Nielsen (2005). Molecular Signatures of Natural Selection. *Annual Review of Genetics*, vol. 39, pp. 197–218. <https://doi.org/10.1146/annurev.genet.39.073003.112420>. DOI: 10.1146/annurev.genet.39.073003.112420.
- [92] Tomoko Ohta (1973). Slightly Deleterious Mutant Substitutions in Evolution. *Nature*, vol. 246, pp. 96–98. <https://www.nature.com/articles/246096a0>. DOI: 10.1038/246096a0.
- [93] Tomoko Ohta (1992). The Nearly Neutral Theory of Molecular Evolution. *Annual Review of Ecology and Systematics*, vol. 23, pp. 263–286. <https://doi.org/10.1146/annurev.es.23.110192.001403>. DOI: 10.1146/annurev.es.23.110192.001403.

- [94] G. Steven Olley, Ariel Pakes (1996). The Dynamics of Productivity in the Telecommunications Equipment Industry. *Econometrica*, Vol. 64, No. 6, pp. 1263-1297 (circulated as NBER Working Paper 3977). <https://www.nber.org/papers/w3977>. DOI: 10.2307/2171831.
- [95] Richard Perline (2005). Strong, Weak and False Inverse Power Laws. *Statistical Science*, vol. 20, no. 1. <https://projecteuclid.org/journals/statistical-science/volume-20/issue-1/Strong-Weak-and-False-Inverse-Power-Laws/10.1214/088342304000000215.full>. DOI: 10.1214/088342304000000215.
- [96] George R. Price (1970). Selection and Covariance. *Nature*, Vol. 227, pp. 520-521. https://www.dynamics.org/Altenberg/LIBRARY/REPRINTS/Price_selection_Nature.1970.pdf. DOI: 10.1038/227520a0.
- [97] William J. Reed (2001). The Pareto, Zipf and other power laws. *Economics Letters*, vol. 74, no. 1, pp. 15-19. <https://www.sciencedirect.com/science/article/pii/S0165176501005249>. DOI: 10.1016/S0165-1765(01)00524-9.
- [98] James Rosindell, Stephen P. Hubbell, Rampal S. Etienne (2011). The Unified Neutral Theory of Biodiversity and Biogeography at Age Ten. *Trends in Ecology and Evolution*, vol. 26, no. 7, pp. 340-348. <https://pubmed.ncbi.nlm.nih.gov/21561679/>. DOI: 10.1016/j.tree.2011.03.024.
- [99] Esteban Rossi-Hansberg, Mark L. J. Wright (2007). Establishment Size Dynamics in the Aggregate Economy. *American Economic Review*, vol. 97, no. 5, pp. 1639-1666. <https://ideas.repec.org/a/aea/aecrev/v97y2007i5p1639-1666.html>. DOI: 10.1257/aer.97.5.1639.
- [100] Alexander Saichev, Yannick Malevergne, Didier Sornette (2010). Theory of Zipf's Law and Beyond. *Springer, Lecture Notes in Economics and Mathematical Systems (book)*. <https://link.springer.com/book/10.1007/978-3-642-02946-2>. DOI: 10.1007/978-3-642-02946-2.
- [101] Herbert A. Simon (1955). On a Class of Skew Distribution Functions. *Biometrika*, vol. 42, no. 3/4. <https://snap.stanford.edu/class/cs224w-readings/Simon55Skewdistribution.pdf>. DOI: 10.2307/2333389.
- [102] Herbert A. Simon, Charles P. Bonini (1958). The Size Distribution of Business Firms. *American Economic Review*, vol. 48, no. 4. <https://www.jstor.org/stable/1808270>.
- [103] Michael H. R. Stanley, Luís A. N. Amaral, Sergey V. Buldyrev, Shlomo Havlin, Heiko Leschhorn, Philipp Maass, Michael A. Salinger, H. Eugene Stanley (1996). Scaling Behaviour in the Growth of Companies. *Nature*, vol. 379. <https://amaral.northwestern.edu/publications/scaling-behaviour-in-the-growth-of-companies/>. DOI: 10.1038/379804a0.
- [104] Josef Steindl (1965). Random Processes and the Growth of Firms: A Study of the Pareto Law. *Charles Griffin & Co. (London)*. <https://www.amazon.com/Random-Processes-Growth-Firms-Pareto/dp/0852640633>.

- [105] Michael P. H. Stumpf, Mason A. Porter (2012). Critical Truths About Power Laws. *Science*, vol. 335, issue 6069, pp. 665–666. <https://www.science.org/doi/10.1126/science.1216142>. DOI: 10.1126/science.1216142.
- [106] John Sutton (1997). Gibrat’s Legacy. *Journal of Economic Literature*, vol. 35, no. 1. <https://www.jstor.org/stable/2729692>.
- [107] Chad Syverson (2011). What Determines Productivity? *Journal of Economic Literature*, Vol. 49, No. 2, pp. 326–365. <https://www.aeaweb.org/articles?id=10.1257/jel.49.2.326>. DOI: 10.1257/jel.49.2.326.
- [108] Fumio Tajima (1989). Statistical Method for Testing the Neutral Mutation Hypothesis by DNA Polymorphism. *Genetics*, vol. 123, no. 3, pp. 585–595. <https://pmc.ncbi.nlm.nih.gov/articles/PMC1203831/>. DOI: 10.1093/genetics/123.3.585.
- [109] Igor Volkov, Jayanth R. Banavar, Stephen P. Hubbell, Amos Maritan (2003). Neutral theory and relative species abundance in ecology. *Nature*, vol. 424, pp. 1035–1037. <https://doi.org/10.1038/nature01883>. DOI: 10.1038/nature01883.
- [110] G. A. Watterson (1975). On the Number of Segregating Sites in Genetical Models Without Recombination. *Theoretical Population Biology*, vol. 7, pp. 256–276. [https://doi.org/10.1016/0040-5809\(75\)90020-9](https://doi.org/10.1016/0040-5809(75)90020-9). DOI: 10.1016/0040-5809(75)90020-9.
- [111] G. Udny Yule (1925). A Mathematical Theory of Evolution, Based on the Conclusions of Dr. J. C. Willis, F.R.S. *Philosophical Transactions of the Royal Society of London, Series B*, vol. 213. <https://royalsocietypublishing.org/doi/10.1098/rstb.1925.0002>. DOI: 10.1098/rstb.1925.0002.