

Economics

A redevelopment of economic theory without parallel universes

Ole Peters and Alexander Adamou



2019/11/27 at 20:14:37

Contents

1	Why this book?	4
1.1	The game	5
1.1.1	Averaging over many trials	6
1.1.2	Averaging over time	8
I	Tools	11
2	Tools	12
2.1	Random variable	13
2.2	Expectation value	15
2.3	Ensemble average	15
2.4	Stochastic process	17
2.5	Time average	19
2.6	The game, revisited	20
2.7	Ergodicity	26
2.8	Changes and stability	28
2.9	Normal distribution	29
2.10	Brownian motion	29
2.11	Geometric Brownian motion	33
2.12	Itô calculus	36
II	Microeconomics	38
3	Decisions in a riskless world	39
3.1	Models and science fiction	40
3.2	The decision axiom	40
3.3	Growth rates	41
3.3.1	Additive growth rate	41
3.3.2	Exponential growth rate	42
3.3.3	General growth rate	43
3.4	Decisions in a deterministic world	45
3.4.1	Different magnitudes	45
3.4.2	Different magnitudes and times: discounting	46

4	Decisions in a risky world	50
4.1	Perturbing the process	50
4.1.1	Perturbed additive dynamics	51
4.1.2	Perturbed multiplicative dynamics	52
4.1.3	Perturbed general dynamics	53
4.2	The appropriate growth rate is ergodic	54
4.3	Ergodicity economics decision algorithm	55
4.4	Relation to earlier economic theories	56
4.4.1	Gamble: random number and duration	56
4.4.2	Repeated gamble: towards a wealth process	59
4.4.3	Expected wealth and expected utility	61
4.5	From ergodicity transformation to dynamic and back – Itô	65
4.5.1	Itô setup	65
4.5.2	From ergodicity transformation to wealth process	65
4.5.3	From wealth process to ergodicity transformation	68
5	Decisions in the real world	72
5.1	The St. Petersburg paradox	72
5.2	The Copenhagen experiment	75
5.3	Insurance	79
5.4	Setup: a shipping insurance contract	80
5.4.1	Ergodicity economics solution	82
5.4.2	The classical solution of the insurance puzzle	84
III	Macroeconomics	85
6	People	86
6.1	Every man for himself	87
6.1.1	Log-normal distribution	87
6.1.2	Two growth rates	88
6.1.3	Measuring inequality	89
6.1.4	Wealth condensation	90
6.1.5	Rescaled wealth	91
6.1.6	u -normal distributions and Jensen’s inequality	92
6.1.7	power-law resemblance	93
6.2	Finite populations	94
6.2.1	Sums of log-normals	94
6.2.2	The random energy model	98
7	Interactions	102
7.1	Cooperation	103
7.2	Reallocation	107
7.2.1	Introduction	107
7.2.2	The ergodic hypothesis in economics	107
7.2.3	Reallocating geometric Brownian motion	109
7.2.4	Model behaviour	110
7.2.5	Derivation of the stable distribution	112
7.2.6	Moments and convergence times	114
7.2.7	Fitting United States wealth data	116

8	Markets	119
8.1	Optimal leverage	120
8.1.1	A model market	120
8.1.2	Leverage	121
8.1.3	Portfolio theory	121
8.1.4	Sharpe ratio	124
8.1.5	Expected return	124
8.1.6	Growth rate maximisation	125
8.2	Stochastic market efficiency	127
8.2.1	A fundamental measure	127
8.2.2	Relaxing the model	127
8.2.3	Efficiency	128
8.2.4	Stability	129
8.2.5	Prediction accuracy	131
8.3	Applications of stochastic market efficiency	131
8.3.1	A theory of noise – Sisyphus discovers prices	132
8.3.2	Solution of the equity premium puzzle	132
8.3.3	Central-bank interest rates	134
8.3.4	Fraud detection	135
8.4	Real markets – stochastic market efficiency in action	135
8.4.1	Backtests of constant-leverage investments	136
8.4.2	Data sets	138
8.4.3	Full time series	139
8.4.4	Shorter time scales	140
8.5	Discussion	142
	Acronyms	142
	List of Symbols	143
	References	146

Chapter 1

Why this book?

{chapter:Why}

In this introductory chapter we lay the conceptual foundation for the rest of what we have to say about redeveloping economic theory. In Sec. 1.1 we play a simple coin-toss game and analyze it numerically, by Monte Carlo simulation, and analytically, with pen and paper. The game motivates the introduction of the expectation value and the time average, which in turn lead to a discussion of ergodic properties.

1.1 The game

{section:The_game}

Imagine we offer you the following game: we toss a coin, and if it comes up heads we increase your monetary wealth by 50%; if it comes up tails we reduce your wealth by 40%. We're not only doing this once, we will do it many times, for example once per week for the rest of your life. Would you accept the rules of our game? Would you submit your wealth to the dynamic our game will impose on it?

Your answer to this question is up to you and will be influenced by many factors, such as the importance you attach to wealth that can be measured in monetary terms, whether you like the thrill of gambling, your religion and moral convictions and so on.

In these notes we will mostly ignore these factors. We will build an extremely simple model of your wealth, which will lead to an extremely simple and powerful model of the way you make decisions that affect your wealth. We are interested in analyzing the game mathematically, which requires a translation of the game into mathematics. We choose the following translation: we introduce the key variable, $x(t)$, which we refer to as “wealth”. We refer to t as “time”. It should be kept in mind that “wealth” and “time” are just names that we've given to mathematical objects. We have chosen these names because we want to compare the behaviour of the mathematical objects to the behaviour of wealth over time, but we emphasize that we're building a model – whether we write $x(t)$, or $\text{wealth}(\text{time})$, or $\odot(\frac{t}{2})$ makes no difference to the mathematics.

The usefulness of our model will be different in different circumstances, ranging from completely meaningless to very helpful. There is no substitute for careful consideration of any given situation, and labelling mathematical objects in one way or another is certainly none.

Having got these words of warning out of the way, we define our model¹ of the dynamics of your wealth under the rules we just specified. At regular intervals of duration δt we randomly generate a factor r with each possible value $r_i \in \{0.6, 1.5\}$ occurring with probability 1/2,

$$r = \begin{cases} 0.6 & \text{with probability } 1/2 \\ 1.5 & \text{with probability } 1/2 \end{cases} \quad (1.1) \quad \{\text{eq:law}\}$$

and multiply current wealth by that factor, so that

$$x(t + \delta t) = x(t)r. \quad (1.2) \quad \{\text{eq:gamble}\}$$

We have good reasons to suspect that this model will do something interesting. It's not just a silly game because it has one important property: it's multiplicative. Every time the coin is tossed, wealth is *multiplied* by one of the two possible factors. This introduces a reference point, or state-dependence, of the absolute size of the change. Processes with this property are very common in nature – the amount of anything that reproduces changes in proportion to what's already there. In the absence of other constraints, such as limited resources, this applies to the dynamics of the biomass of a cell culture. In fact, life itself has been defined as that which produces more of itself [38]: life and evolution begin with the minimal chemical structure that can copy itself. The rest,

¹For those in the know: “our” coin toss is a discrete version of geometric Brownian motion, the workhorse model of financial mathematics and much more.

in a sense, is details. We will see, especially in Chap. 7, how multiplicativity generates interesting structure in – including but not limited to – economics.

Without discussing in depth how realistic a representation of your wealth this model is (for instance your non-gambling related income and spending are not represented in the model), and without discussing whether randomness truly exists and what the meaning of a probability is, we simply switch on a computer and simulate what might happen. You may have many good ideas of how to analyze our game with pen and paper, but we will just generate possible trajectories of your wealth and pretend we know nothing about mathematics or economic theory. Figure 1.1 is a trajectory of your wealth, according to our computer model as it might evolve over the course of 52 time steps (corresponding to one year given our original setup).

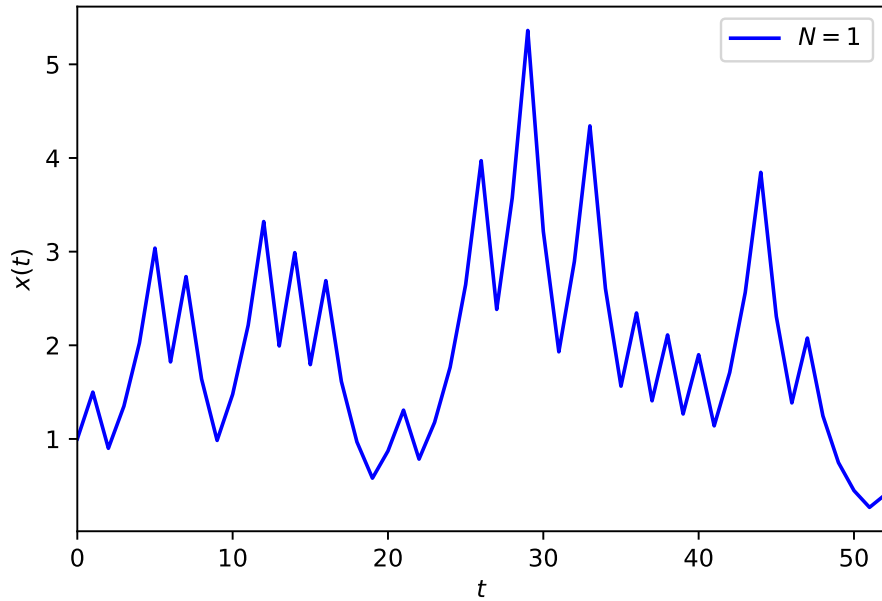


Figure 1.1: Wealth $x(t)$ resulting from a computer simulation of our game, as specified by (Eq. 1.1) and (Eq. 1.2), for 52 time steps (corresponding to one year in the given setup).

{fig:1_1}

A cursory glance at the trajectory does not reveal much structure. Of course there are regularities, for instance at each time step $x(t)$ changes, but no trend is discernible – does this trajectory have a tendency to go up, does it have a tendency to go down? Neither? What are we to learn from this simulation? Perhaps we conclude that playing the game for a year is quite risky, but is the risk worth taking?

1.1.1 Averaging over many trials

The trajectory in Fig. 1.1 doesn't tell us much about overall tendencies. There is too much noise to discern a clear signal. A common strategy for getting rid of noise is to try again. And then try again and again, and look at what happens on average. An example of the technique is Shannon's error-correcting code:

instead of sending the message 0 (or 1), send the relevant digit 3 times. The recipient averages over the received digits and takes the closest possibility. If one out of three digits was miscommunicated because of noise, the code nonetheless recovers the original message: averaging gets rid of noise.

So let's try this in our case and see if we can make sense of the game. In Fig. 1.2 we average over a finite number, N , of trajectories. We call a collection of trajectories an ensemble. We shall see that in the limit $N \rightarrow \infty$ the ensemble average converges to the expectation value, and indeed the terms “ensemble average” and “expectation value” are synonyms. To avoid confusion we will be explicit when N is finite: in Fig. 1.2 we plot “finite-ensemble averages.”

DEFINITION: Finite-ensemble average

The finite-ensemble average of the quantity z at a given time t is

$$\langle z(t) \rangle_N = \frac{1}{N} \sum_{i=1}^N z_i(t), \quad (1.3) \quad \{\text{eq:f_ens}\}$$

where i indexes a particular realization of $z(t)$ and N is the number of realizations included in the average.

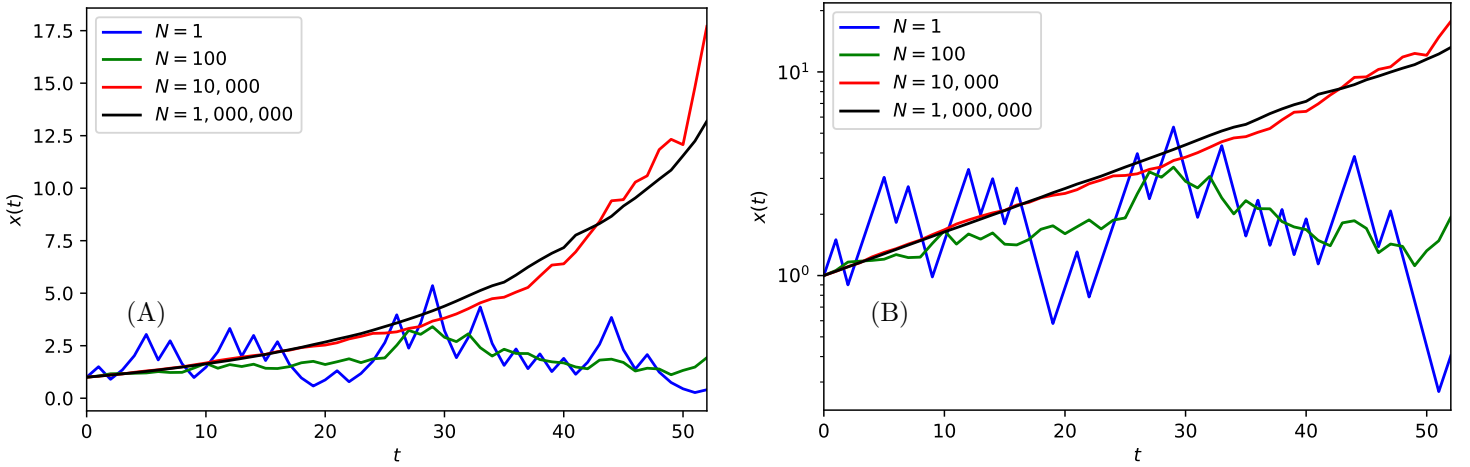


Figure 1.2: Finite-ensemble averages $\langle x(t) \rangle_N$ for ensemble sizes $N = 1, 10^2, 10^4, 10^6$. The noise in the finite-ensemble average diminishes systematically as N increases. (A) on linear scales the multiplicative (non-linear) nature of the process is apparent, (B) on logarithmic scales the multiplicative process is additive in time, and the finite-ensemble average for $N = 10^6$ is a straight line except for small fluctuations.

{fig:1_2}

As expected, the more trajectories are included in the average, the smaller the fluctuations of that average. For $N = 10^6$ hardly any fluctuations are visible. Since the noise-free trajectory points up it is tempting to conclude that

my own wealth will similarly go up and conclude that the risk of the game is worth taking. This reasoning has dominated economic theory for about 350 years now. But it is flawed. The correction of this flaw and its far-reaching consequences constitute our research program.

1.1.2 Averaging over time

Does our analysis necessitate the conclusion that the gamble is worth taking? Of course it doesn't, otherwise we wouldn't be belabouring this point. Our critique will focus on the type of averaging we have applied – we didn't play the game many times in a row as would correspond to the real-world situation of repeating the game once a week for the rest of your life. Instead we played the game many times in parallel, which corresponds to a different setup².

We therefore try a different analysis. Figure 1.3 shows another simulation of your wealth. This time we don't show an average over many trajectories but a simulation of a single trajectory over a long time. Noise is removed also in this case but in a different way: to capture visually what happens over a long time we have to zoom out – more time has to be represented by the same amount of space on the page. In the process of this zooming-out, small short-time fluctuations will be diminished. Eventually the noise will be removed from the system just by the passage of time.

²This different setup could be realised by splitting your wealth into N equal parts and betting each part in a different sequence of independent coin tosses. But the rules of our game as we defined it don't allow that. Another interesting setup allows the gambler to choose what proportion of his wealth he wants to wager. These conditions were studied by Kelly [24] and are known to every professional poker player. In the present setup, betting $1/4$ of your wealth in each coin toss will lead to the fastest possible growth if you're allowed to choose your wager. You can think of the proportion not wagered as frozen in time: it ensures that, to some extent, you can restore the conditions before the coin toss, which is a bit like allowing you to step back in time.

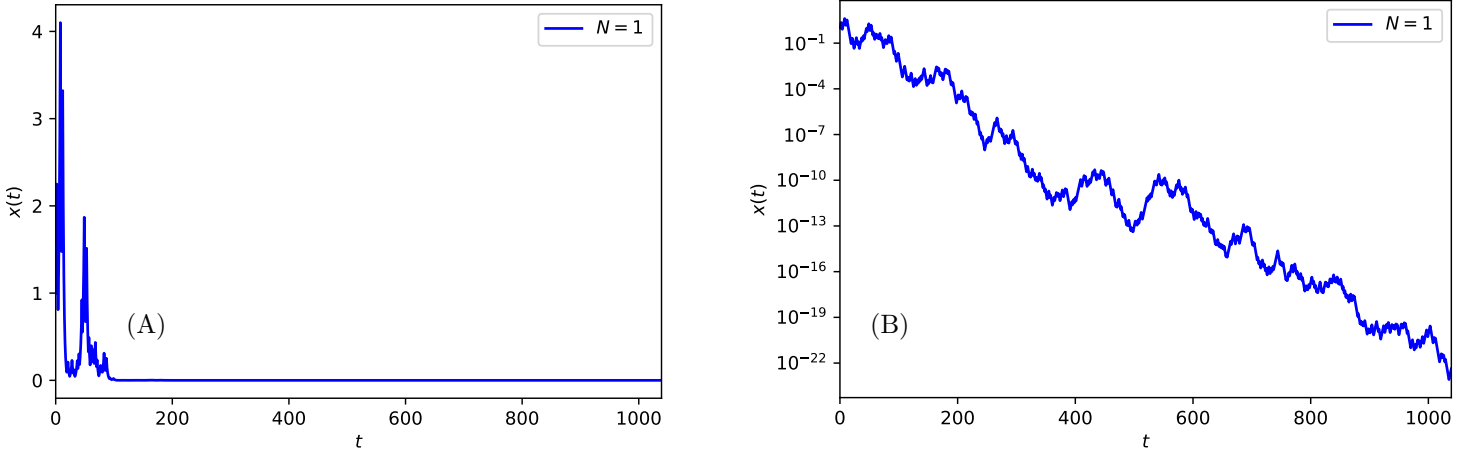


Figure 1.3: Single trajectory over 1,040 time steps, corresponding to 20 years in our setup. (A) on linear scales all we see is wealth quickly dropping to zero, (B) logarithmic scales are more appropriate for the multiplicative process and reveal systematic exponential decay.

{fig:1_3}

The trajectory in Fig. 1.3 is, of course, random, but the apparent trend emerging from the randomness strongly suggests that our initial analysis in Fig. 1.2 does not reflect what happens over time in a single system. Several important messages can be derived from the observation that an individual trajectory grows more slowly (or decays faster) over time than an average of a large ensemble.

1. An individual whose wealth follows (Eq. 1.2) will make poor decisions if he uses the finite-ensemble average of wealth as an indication of what is likely to happen to his own wealth.
2. The performance of the average (or aggregate) wealth of a large group of individuals differs systematically from the performance of an individual's wealth. In our case large-group wealth grows (think [gross domestic product \(GDP\)](#)), whereas individual wealth decays.
3. For point 2 to be possible, *i.e.* for the average to outperform the typical individual, wealth must become increasingly concentrated in a few extremely rich individuals. The wealth of the richest individuals must be so large that the average becomes dominated by it, so that the average can grow although almost everyone's wealth decays. Inequality increases in our system.

The two methods we've used to eliminate the noise from our system are well known. The first method is closely related to the mathematical object called the “expectation value,” and the second is closely related to the object called the “time average.”

Summary of Chap. 1

Suppressed.

Part I

Tools

Chapter 2

Tools

{chapter:Tools}

The ergodicity question – whether time averages are identical to expectation values – is the key to our redevelopment of formal economics. This is because ergodicity hadn't been established as a concept when the original formalism was developed. The scientific search for stable structures leads to constants in deterministic settings. When randomness is introduced, the role previously played by constants is taken on by ergodic observables. We also introduce the concepts of a random variable, a stochastic process, scalars as representations of transitive preferences, logarithms and exponentials, and dimensional analysis.

In Sec. 2.10 we notice that wealth on logarithmic scales follows a random walk in our game, and we relate this to Brownian motion, as the continuous-time limit of the random walk. This allows us to introduce Brownian motion and its scaling properties that are robust enough to yield insights into more complicated models.

Finally, we ask in Sec. 2.11 what wealth in our game is doing in the continuum limit but on linear scales. This takes us to geometric Brownian motion, which will be our starting point for much of the rest of these lectures. We derive ensemble-average and time-average growth rates for geometric Brownian motion, by explicitly taking the continuous-time limit, and then state the key result of Itô calculus, (Eq. 2.72) and (Eq. 2.73), which allows an easier derivation of these growth rates and will be relied on in later chapters.

Some historical perspective is provided to understand the prevalence or absence of key concepts in modern economic theory and other fields. The emphasis is on introducing concepts and useful machinery, with applications in later chapters.

2.1 Random variable

{section:random_variable}

In economics, as elsewhere, we are often interested in ‘experiments’ whose outcomes we don’t yet know. Examples are each coin toss in the game in Sec. 1.1 and the result of a football match. We might know something about the experiment, such as the possible outcomes and that some are more plausible than others, but we are ignorant of the actual outcome. Luckily, we can use probability theory, a branch of mathematics, to build models of our ignorance.

Often the experimental outcomes have, or can be mapped to, numerical values. For example, each coin toss in the game has possible outcomes heads and tails, which correspond to wealth multipliers 1.5 and 0.6. In such cases, we model the uncertain numerical value as a *random variable*. We assume you have seen random variables before and we will not give a lengthy technical account. Instead, we recommend the two-page discussion in [61, p. 2], whose key points we reproduce here.

A random variable, Z , is defined by:

- the set of its possible values, $\{z_j\}$; and
- a probability distribution, $P[\]$, over this set.

The set of possible values of Z may be continuous, like the interval $(3, 12)$ or the real numbers, $\mathbb{R}$; discrete, like $\{4, 7.8, 29\}$ or the integers, $\mathbb{Z}$; or a combination of the two. The probability distribution is a function which maps values to probabilities. We define an event as Z taking the value z , which we write as $Z = z$. This is strictly an abuse of notation because Z is a random variable, and z is a scalar (so they can’t be identical as the “=” sign suggests). But this is common notation, and we will use it here. Similarly, we may write $Z < a$ to denote the union of all events where Z takes a value $z < a$ etc. In terms of probabilities, we write

$$P[Z = z] = p \quad (2.1)$$

to mean that the event $Z = z$ is associated with the probability p . A specific value, z , is sometimes called an ‘outcome’, an ‘instance’ or a ‘realisation’ of the random variable, Z . While not obligatory, it is a common convention to denote random variables in upper case and realisations in lower case. Later in these lecture notes, where this won’t lead to confusion, we will often overload notation and use lower case letters for both random variables and realisations. This is a little less precise but easier on the eye.

Forget, for the moment, what probabilities might mean in the context of an experiment. In purely mathematical terms, they are just real numbers associated with outcomes. The probability distribution has two constraints:

- the probability of any outcome must be non-negative; and
- the probability of an outcome that is certain to happen, *i.e.* that includes all possible outcomes, must be one.

The latter is a normalisation condition which fixes the scale of the probabilities.

Discrete random variable

When the set of outcomes is discrete, say $\{z_1, \dots, z_M\}$, we assign a probability, p_j , to each outcome, z_j , such that

$$\mathbb{P}[Z = z] = \begin{cases} p_j & \text{if } z = z_j \\ 0 & \text{otherwise} \end{cases} \quad (2.2)$$

and

$$\sum_{j=1}^M p_j = 1. \quad (2.3)$$

Continuous random variable

Most of the models we will study contain random variables whose possible outcomes form a continuous set. In this case, Z can take uncountably many values, to which we can't assign non-zero probabilities while maintaining the normalisation condition. Instead, we assign probabilities to intervals. We define a **probability density function (PDF)**, $\mathcal{P}_Z(z)$, such that the probability of Z being in the interval (a, b) is given by

$$\mathbb{P}[a \leq Z \leq b] = \int_a^b \mathcal{P}_Z(z) dz. \quad (2.4)$$

The **PDF** is a non-negative function, $\mathcal{P}_Z(z) \geq 0$, normalised so that the probability of the certain outcome, *i.e.* the integral over all possible outcomes, is one:

$$\int_{-\infty}^{+\infty} \mathcal{P}_Z(z) dz = 1. \quad (2.5)$$

Note the difference between subscript and argument: $\mathcal{P}_Z(z)$ is the probability density of the random variable Z at value z . You might find it helpful to think of $\mathcal{P}_Z(z)\delta z$ as the approximate probability of Z being in the small interval $(z, z + \delta z)$ close to z .

Interpretation

So far we have said nothing about the meaning of probabilities: formally, they are simply numbers assigned to outcomes of random variables. Giving these numbers meaning is a separate step.

In these notes, we will use a frequentist interpretation of probabilities, by which we mean the following. Imagine many separate experiments. For example, suppose we toss a coin N times, and record the number of times, n , that a particular outcome occurs. Under the frequentist interpretation, the probability associated with this outcome is its relative frequency, n/N , in the limit $N \rightarrow \infty$. In the coin toss example, the probability assigned to heads would be 0.5 if the coin were unbiased; if biased, it would be some other number between 0 and 1. Formally, in terms of a random variable, the frequentist interpretation means that n/N will approach the probability of n when we generate N instances of the random variable.

Note also that time does not appear in the formal definition of a random variable. The moment we try to illustrate what a random variable might model, we introduce time. For instance by saying “we toss a coin, and we model its *future* state (heads or tails) as a random variable.” But the mathematical object

that is the random variable doesn't know about this interpretation. It just knows about possible values, associated with numbers between zero and one.

Nor do random variables usually depend on time. Of course, there is nothing stopping us from using a probability distribution which does depend on time or, indeed, any other variable, like the day of the week or the country we are in. Such parameterisations of the random variable do not change fundamentally its mathematical structure: a set of outcomes and associated probabilities. When we consider probability distributions of random variables that do depend on time, as we do for wealth, $x(t)$, in the coin tossing game, we will make the time dependence explicit. By default we assume random variables are time-independent

2.2 Expectation value

The *expectation value*¹ is a property of a random variable. Denoted by $E[Z]$, it is the probability-weighted average of the realisations of Z .

DEFINITION: **Expectation value**

If Z is a discrete random variable, its expectation value is the sum of the possible realisations, z_j , weighted by their probabilities, p_j :

$$E[Z] = \sum_j p_j z_j. \quad (2.6) \quad \{\text{eq:exp_sum}\}$$

If Z is a continuous random variable, its expectation value is the integral over the possible realisations weighted by the probability density:

$$E[Z] = \int_{-\infty}^{+\infty} \mathcal{P}_Z(z) z dz. \quad (2.7)$$

The sum or the integral may not exist, in which case the random variable does not have an expectation value. One examples of this are the payout of the St Petersburg lottery, which we will encounter in Sec. 5.1. Another example is any fat-tailed random variable with PDF of the form:

$$\mathcal{P}_Z(z) = \begin{cases} (\tau - 1)z^{-\tau} & z \geq 1 \\ 0 & z < 1 \end{cases} \quad (2.8)$$

for $\tau \leq 2$

2.3 Ensemble average

The *ensemble average* is conceptually different from the expectation value, although we will see that it always is identical in value. Instead of weighting the average of the possible realisations of Z by their probabilities, we take an unweighted average of a collection of generated realisations. For a finite number of realisations, we call this the *finite-ensemble average* and denote it by $\langle Z \rangle_N$.

¹Also known as expected value, mathematical expectation, first moment, and mean.

DEFINITION: Finite-ensemble average

The finite-ensemble average of a random variable, Z , is the average over a finite number, N , of generated realisations,

$$\langle Z \rangle_N \equiv \frac{1}{N} \sum_{i=1}^N z_i, \quad (2.9) \quad \{\text{eq:f_ens}\}$$

where z_i denotes the i^{th} realisation of Z .

You may know this as the sample average or sample mean. Note that it is an additive average, meaning that $\langle Z \rangle_N$ is the value by which each of the N realisations of Z must be replaced if we want them all to be equal without altering their *sum*.

The finite-ensemble average is itself a random variable: each finite ensemble of size N can contain different realisations of Z which sum to a different total. As $N \rightarrow \infty$, this random average may converge with probability one to a unique constant. If it does, we call this limiting value the ensemble average and denote it by $\langle Z \rangle$.

DEFINITION: Ensemble average

The ensemble average of a random variable, Z , is the $N \rightarrow \infty$ limit of the finite-ensemble average,

$$\langle Z \rangle \equiv \lim_{N \rightarrow \infty} \langle Z \rangle_N = \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N z_i, \quad (2.10) \quad \{\text{eq:ens}\}$$

where z_i denotes the i^{th} realisation of Z .

The limit is not guaranteed to exist, in which case Z has no ensemble average, although finite-ensemble averages can always be computed. We will sometimes refer to the imagined infinite collection of realisations as the “ensemble” of Z .

Incidentally, if $\langle Z \rangle$ is divergent, then $\langle Z \rangle_N$ will grow systematically with N . If $\langle Z \rangle$ represents an average taken over N successive measurements (a time series), then a fat-tailed Z is often as good a model as a non-stationary (time-dependent) Z .

Note that, unlike the expectation value, the ensemble average of a random variable is not defined in terms of probabilities. It is the quantity to which averages over realisations converge as the number of realisations grows. Recall that we followed this procedure when analysing the coin game. At each round of the gamble, we found finite-ensemble averages of simulated wealth for increasingly large samples, from one to one million, plotted against time in Fig. 1.2.

It is laborious to compute averages of ever larger finite ensembles, in the hope of discerning their convergence to a limit. Fortunately, we can show, under the frequentist interpretation of probability introduced in Sec. 2.1, that the ensemble average of a random variable is equal to its expectation value. We do this here for the discrete case, noting that a similar proof can be offered for the continuous case.

Proof. Denote by n_j the number of times z_j is observed in N realisations of Z . The finite-ensemble average can be written as

$$\langle Z \rangle_N = \frac{1}{N} \sum_{i=1}^N z_i = \sum_j \frac{n_j}{N} z_j, \quad (2.11)$$

where the subscript i indexes a particular realisation of z and the subscript j indexes a possible value of z . Circularly – because we work with frequentist probabilities – the relative frequency of each possible value, n_j/N , converges almost surely in the limit $N \rightarrow \infty$ to its (frequentist) probability, p_j . Thus, we find

$$\lim_{N \rightarrow \infty} \langle Z \rangle_N = \sum_j p_j z_j = \mathbb{E}[Z]. \quad (2.12)$$

□

This result, commonly known as the *law of large numbers*, is exceedingly useful. It means that we no longer need many realisations of a random variable to compute its ensemble average. Instead, if we know the probability distribution, we can compute its expectation value straightforwardly as a weighted average or integral. It also means that, from now on, we can use ensemble average and expectation value interchangeably.

History: The invention of the expectation value

Coming soon.

2.4 Stochastic process

In the coin game, the wealth multiplier, r , at each round is a random variable with outcomes $\{0.6, 1.5\}$ and probabilities $\{0.5, 0.5\}$. Starting from $x(0) = \$1$, successive realisations of the wealth updating rule,

$$x(t + \delta t) = x(t)r, \quad (2.13) \quad \{\text{eq:x_update}\}$$

generate a function of time, $x(t)$, which describes one possible evolution of wealth in the game. If we start again, generating fresh realisations of r at each time step, we get a second function of time, almost certain to be different from the first. If run the game N times, we get a set of wealth evolution functions, $\{x_1(t), \dots, x_N(t)\}$.

This situation reminds us of a random variable, except that each realisation is now a function (x) of a parameter (t). Each $x_j(t)$ is an ordered set of real numbers (one for each t), rather than a single real number. We call such functions *trajectories*. Many natural phenomena are like the coin game, in that they result not in single observations but in ordered sequences of connected observations, about whose values we are uncertain. To model such phenomena, we need a new type of mathematical object, the *stochastic process*, denoted $Y(t)$. We define this analogously to the random variable, as:

- a set of possible trajectories, $\{y(t)\}$; and
- a probability distribution over this set.

More technical presentations are available and, again, we point you to [61, p. 52] for greater depth. The set of trajectories may be countable and associated with a discrete set of probabilities, or uncountable and associated with a PDF. We will always interpret the parameter, t , as time.

This might seem like a lot to take in, but conceptually it's fairly simple. A stochastic process is nothing more than a family of trajectories from which we can select at random, according to a probability distribution. It is the generalisation of the random variable from single numbers to functions of time. If you are already familiar with stochastic processes, this may not be how you think of them. We illustrate our perspective in Fig. 2.1.

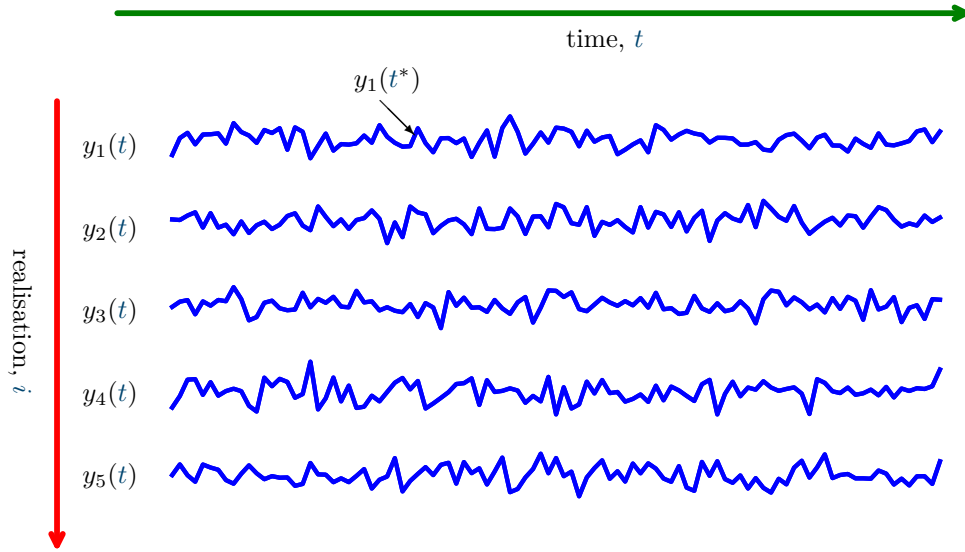


Figure 2.1: Realisations of a stochastic process, $Y(t)$. Each trajectory, $y_i(t)$ is a function of time, which runs horizontally across the page. Fixing time, $t = t^*$, results in a random variable, $Y(t^*)$, one realisation of which, $y_1(t^*)$, is shown.

{fig:sp_grid}

This picture suggests another way of viewing the stochastic process: if we choose a time, $t = t^*$, then $Y(t^*)$ is a random variable, whose possible realisations form the set of possible function values, $\{y(t^*)\}$. So, if you prefer, you can think of a stochastic process as a family of random variables, with probability distributions parametrised by time.

Generating stochastic processes

Switching to practical considerations, there are two main ways of generating a trajectory of a stochastic processes, say on a computer. The first is consistent with our conceptualisation of selecting one from a family of trajectories. For example, if the possible trajectories were finitely many and equally probable, then we could label them with integers and draw one from a discrete uniform distribution. All the uncertainty is resolved at the start, in a single realisation of a random variable.

The other approach is to generate a trajectory sequentially, by successive

application of the rule that tells us which possible values of $y_i(t + \delta t)$ are compatible with a known value of $y_i(t)$. This is how we generate wealth trajectories in the coin game, where the rule is to multiply current wealth by the realisation of a random variable. Realising a random variable at every time step sounds laborious. However, in many models of natural phenomena, it is easier to write down a realistic updating rule than an expression for all possible trajectories. This is because it is often about the changes of quantities that we can reason ‘physically’. Moreover, many stochastic processes have uncountably many trajectories, making the single realisation approach difficult, because it requires high numerical precision, *i.e.* very long numbers, to be useful.

Therefore, typically we generate trajectories sequentially, using updating rules which we call the *dynamics* of the stochastic process. Usually, where it creates no ambiguity, we don’t distinguish between the stochastic process and a single realisation of it, using lower case to denote both, *e.g.* $x(t)$ for wealth. When it’s important to identify different trajectories, we include an index, *e.g.* $x_i(t)$. Sometimes we refer to these different realisations as *systems* and to the index as the system number.

2.5 Time average

Stochastic processes are models of uncertain quantities which vary across systems and over time. One example is the wealth of a player of the coin game; another example is the total goals scored in a football match. If we fix time in a stochastic process, then we can take the ensemble average of the resulting random variable,

$$\langle Y(t^*) \rangle = \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N y_i(t^*). \quad (2.14)$$

The introduction of time in the stochastic process makes another type of average possible, conceptually different from the ensemble average. Instead of choosing a time and averaging across systems, we can choose a system and average over time. Let’s start by doing this for a finite time period, Δt , to generate a *finite-time average*.

DEFINITION: Finite-time average

The finite-time average of the trajectory $y(t)$ is the average value of the function from t to $t + \Delta t$,

$$\bar{y}_{\Delta t}(t) \equiv \frac{1}{\Delta t} \int_t^{t+\Delta t} y(s) ds. \quad (2.15) \quad \{\text{eq:t_ave_f}\}$$

If $y(t)$ changes only at discrete times, $\{t + \delta t, \dots, t + T\delta t\}$, where $\Delta t = T\delta t$, then the finite-time average can be written as

$$\bar{y}_{\Delta t}(t) = \frac{1}{T\delta t} \sum_{\tau=1}^T y(t + \tau\delta t). \quad (2.16) \quad \{\text{eq:t_ave_f_disc}\}$$

In the same way that the ensemble average is the many-systems limit of the finite-ensemble average, so the *time average* is the long-time limit of the finite-time average.

DEFINITION: Time average

The time average is the long-time limit of the finite-time average,

$$\bar{y}(t) \equiv \lim_{\Delta t \rightarrow \infty} \bar{y}_{\Delta t}(t), \quad (2.17) \quad \{\text{eq:t_ave}\}$$

where this limit exists.

Note that both the finite-time average and the time average are strictly functions of the time, t , at which the averaging starts. In many of the models we consider, dependence on the initial condition vanishes in the limit of long averaging time. If so, we remove it from the notation and write $\bar{y}$.

The two types of average possible in a stochastic process are illustrated in Fig. 2.2.

2.6 The game, revisited

We have introduced some essential tools for studying models such as that of the coin game, where outcomes are uncertain and vary over time. It makes sense, therefore, to apply these tools to the game to see how they deepen our understanding of it.

We pretended to be mathematically clueless when we ran the simulations in Sec. 1.1, hoping to understand the game simply by looking at many possible outcomes. We can now compute exactly the ensemble average of the stochastic process for wealth, $x(t)$, instead of approximating it numerically by a finite-ensemble average. We start by taking the ensemble average of the wealth updating equation, (Eq. 2.13),

$$\langle x(t + \delta t) \rangle = \langle x(t)r \rangle. \quad (2.18) \quad \{\text{eq:step_1}\}$$

We assume the outcomes of coin tosses are independent of the player's wealth

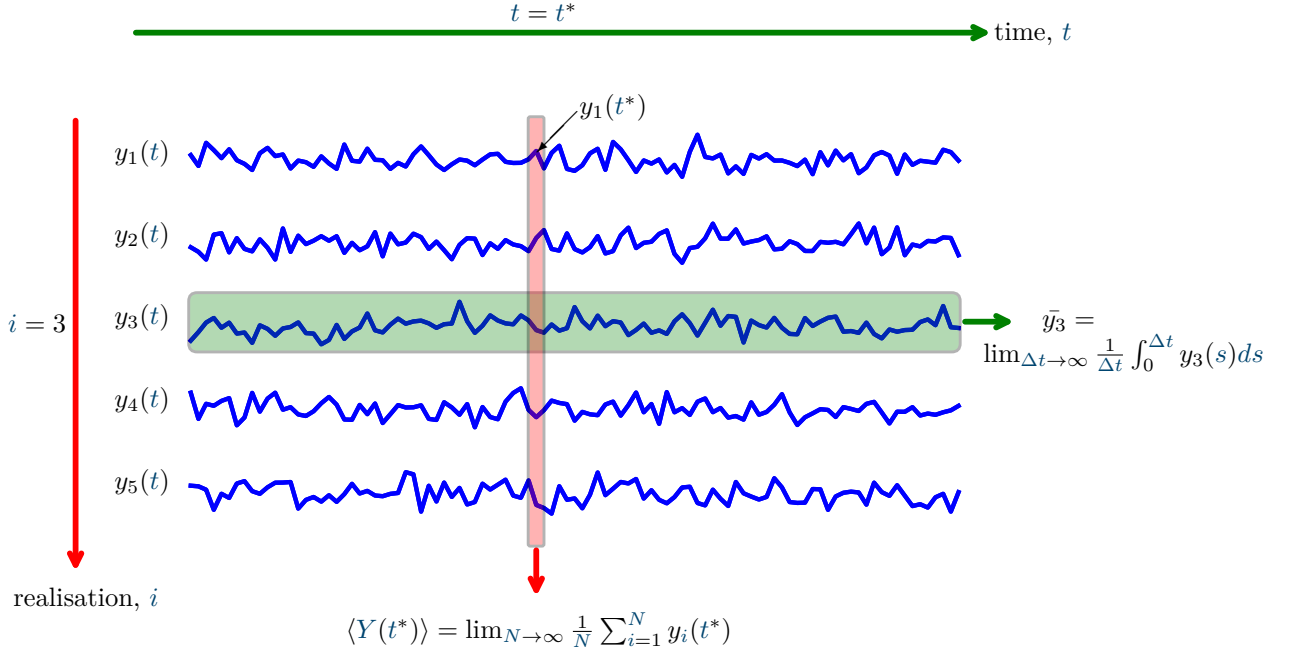


Figure 2.2: Realisations of a stochastic process, $Y(t)$, *cf.* Fig. 2.1. The ensemble average, $\langle Y(t^*) \rangle$, is obtained by choosing a time, $t = t^*$, and averaging over trajectories (red box). The time average is obtained by choosing a trajectory, $i = 3$, and averaging over time (green box).

{fig:ergodic_grid}

and of time, so that $\langle x(t)r \rangle = \langle x(t) \rangle \langle r \rangle$. Thus, (Eq. 2.18) can be written as

$$\langle x(t + \delta t) \rangle = \langle x(t) \rangle \langle r \rangle. \quad (2.19)$$

We can solve this recursively for the wealth after T rounds from a known starting wealth, $x(t_0)$:

$$\langle x(t_0 + T\delta t) \rangle = x(t_0) \langle r \rangle^T. \quad (2.20) \quad \{\text{eq:x_exp_r}\}$$

δt is the duration of a single round of the game, so the total playing time is $\Delta t = T\delta t$.

Recall that r is random variable with outcomes $\{0.6, 1.5\}$ and probabilities $\{0.5, 0.5\}$. The expectation value, $\langle r \rangle$, is, therefore,

$$\langle r \rangle = (0.5)(0.6) + (0.5)(1.5) = 1.05. \quad (2.21)$$

Since this number is greater than one, the ensemble average wealth, $\langle x(t) \rangle$, grows exponentially over time by a factor 1.05 per round of the game.

It is useful to quantify this growth in a standard way, as a growth rate. We will say much more about this in Sec. 3.3, since it is a central concept in ergodicity economics. For the moment, it is enough to know that the appropriate growth rate for a quantity which grows multiplicatively is the coefficient of time, g , in the exponential function, $\exp(gt)$.

We write

$$\langle x(t_0 + \Delta t) \rangle = x(t_0) \exp(g_{\langle} \Delta t) \quad (2.22) \quad \{\text{eq:x_exp_g}\}$$

and equate it to $x(t_0) \langle r \rangle^T$ in (Eq. 2.20) to get the exponential growth rate,

$$g_{\langle} = \frac{T \ln \langle r \rangle}{\Delta t} = \frac{\ln \langle r \rangle}{\delta t}. \quad (2.23)$$

We aren't interested here in whether time is measured in minutes, days, or years, so for simplicity let's just measure it in rounds of the game, *i.e.* $\delta t = 1$ round. This choice gives us a growth rate of $g_{\langle} = \ln(1.05)/1 \approx 4.9\%$ per round.

The positive growth rate of ensemble average wealth is what might have led us to conclude that the game is worth playing, when we analysed it numerically in Sec. 1.1. Figure 2.3 compares the analytical result for the ensemble, derived above, to the numerical results of Fig. 1.2 for finite ensembles. We could now conclude that the case is closed. Mathematics and simulations agree there is some risk involved, but “on average” our wealth grows if we play the game.

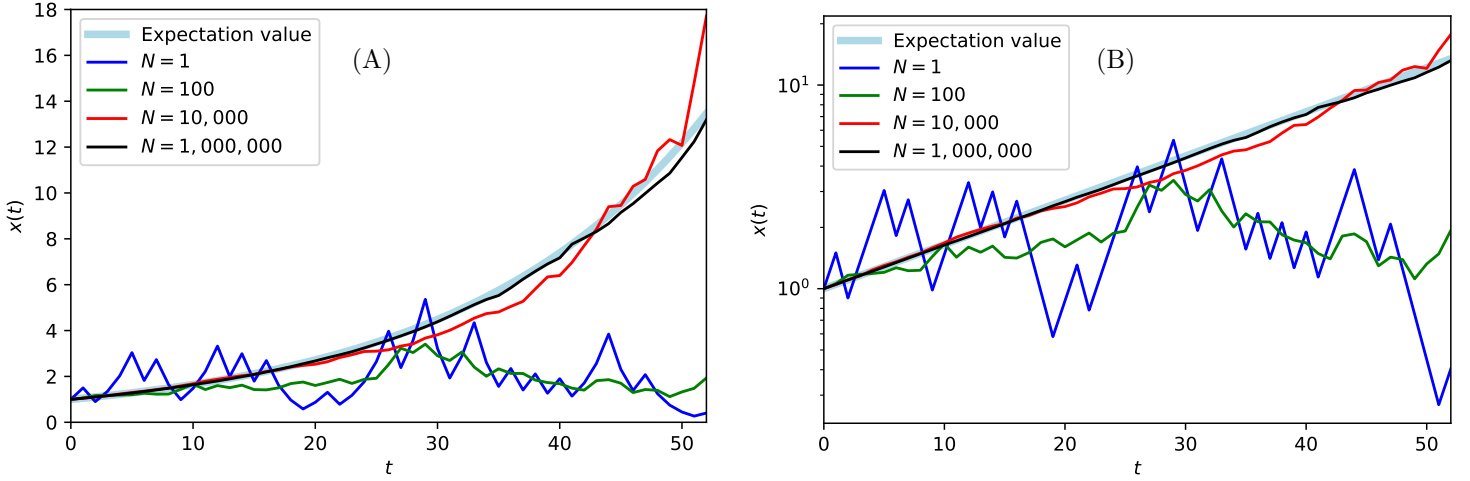


Figure 2.3: Expectation value (thick light blue line) and finite-ensemble averages on (A) linear scales and (B) logarithmic scales. {\text{fig:cf_exp}}

However, we now know that there are two ways in which we can say “on average”: we have shown how wealth averaged over many systems evolves; let's now consider a single wealth trajectory over long time. Starting from $x(t_0)$ and playing for T rounds leads to wealth,

$$x(t_0 + T\Delta t) = x(t_0) \prod_{\tau=1}^T r_{\tau}, \quad (2.24)$$

where $\{r_{\tau}\}$ are random wealth multipliers, all distributed like r , and τ is the index for the round of the game. We have taken no kind of average yet and $x(t_0 + T\Delta t)$ remains a random variable (it is, after all, the outcome of a stochastic process at a specified time).

We follow our previous strategy of expressing this as exponential growth

$$x(t_0 + \Delta t) = x(t_0) \exp(g_T \Delta t), \quad (2.25)$$

where we have labelled the growth rate as g_T to indicate that it is measured after T rounds of the gamble. Equating the two expressions and rearranging for the growth rate, we get

$$g_T = \frac{1}{\Delta t} \ln \left(\prod_{\tau=1}^T r_\tau \right) = \frac{1}{\delta t} \left(\frac{1}{T} \sum_{\tau=1}^T \ln r_\tau \right), \quad (2.26) \quad \{\text{eq:g_x_T}\}$$

where we used the unique property of the logarithm that $\ln(\alpha\beta) = \ln \alpha + \ln \beta$ to convert the product into a sum. Individual wealth grows exponentially at a random growth rate, g_T , whose distribution depends on the number of rounds played, T .

What happens to g_T in the long run? There are two ways of evaluating the $\Delta t \rightarrow \infty$ (or, equivalently, $T \rightarrow \infty$) limit of (Eq. 2.26). Both are illuminating in different ways. First, we can consider the product of wealth multipliers. The effective multiplier which, if applied at each of the T rounds would yield the same product, is

$$r_{\text{eff}} = \left(\prod_{\tau=1}^T r_\tau \right)^{\frac{1}{T}}. \quad (2.27)$$

This is known as the *geometric mean* of $\{r_\tau\}$. Suppose heads appear m times and tails n times. Then

$$r_{\text{eff}} = (1.5)^{\frac{m}{T}} (0.6)^{\frac{n}{T}}, \quad (2.28)$$

which is simply the product of possible realisations of the wealth multiplier raised to their relative frequencies. We know already what happens to relative frequencies the number of rounds grows large: they converge almost surely to their corresponding probabilities, in this case both $1/2$. So, over many rounds, the effective per-round multiplier of wealth converges to

$$r_{\text{time}} = \lim_{T \rightarrow \infty} r_{\text{eff}} = (1.5)^{\frac{1}{2}} (0.6)^{\frac{1}{2}} = (0.9)^{\frac{1}{2}} \approx 0.95. \quad (2.29)$$

This tells a strikingly different story from that told by ensemble average wealth. Instead of growing exponentially by a factor 1.05 per round, a single wealth trajectory over long time decays exponentially by an effective factor of 0.95 per round. The corresponding exponential growth rate, from (Eq. 2.26), is

$$g_{\text{time}} = \frac{\ln r_{\text{time}}}{\delta t} \approx -5.3\% \text{ per round}. \quad (2.30)$$

The alternative route to this result is to take the $T \rightarrow \infty$ limit of the sum in (Eq. 2.26). Something rather interesting – and central to ergodicity economics – happens when we do this:

$$g_{\text{time}} = \lim_{T \rightarrow \infty} \frac{1}{\delta t} \left(\frac{1}{T} \sum_{\tau=1}^T \ln r_\tau \right) = \frac{\langle \ln r \rangle}{\delta t}. \quad (2.31) \quad \{\text{eq:g_time}\}$$

The first equality says that the growth rate of individual wealth over long time is the time average of the logarithm of the wealth multipliers, $\{\ln r_\tau\}$, divided

by the round duration, δt . That’s kind of believable. After all, we are trying to quantify how something is growing over time. The second equality, however, contains a remarkable conceptual insight: we can express this time average as the ensemble average of the random variable, $\ln r$. In other words, provided we know which observable to average over, we can investigate what happens in our model over long time by averaging over many systems!

To maintain sanity, let’s confirm this is the same as the result obtained by finding the effective wealth multiplier:

$$\frac{\langle \ln r \rangle}{\delta t} = \frac{\frac{1}{2} \ln(1.5) + \frac{1}{2} \ln(0.6)}{\delta t} = \frac{\ln \left[(1.5)^{\frac{1}{2}} (0.6)^{\frac{1}{2}} \right]}{\delta t} = \frac{\ln r_{\text{time}}}{\delta t}. \quad (2.32)$$

It checks out.

There are two deep reasons why this works. The first is that the operation we perform on the product of the $\{r_\tau\}$ in (Eq. 2.26) – taking the logarithm and dividing by the duration – produces an arithmetic mean of the $\{\ln r_\tau\}$. In effect, by using the logarithm to transform multiplication to addition, we express the evolution of wealth in the form of a sum, specifically a finite-time average. The second reason is that the $\{\ln r_\tau\}$ are independent and identically distributed (“iid”) random variables, with the same distribution as $\ln r$ and with no dependence on time. It follows that realisations over time and across systems are statistically indistinguishable, so that the time average in (Eq. 2.31) can be replaced by the ensemble average. Another way of seeing this is that the summation index, τ , in (Eq. 2.31) can just as well index different players as different rounds, because of the iid-ness of the $\{r_\tau\}$.

Observables whose time and ensemble averages are equivalent are very special, and we will say much more about them in the next section.

Our computations of g_{time} are consistent with the simulations in Sec. 1.1. Figure 2.4 (B) compares the single trajectory generated in Fig. 1.3 to exponential decay at rate g_{time} . The trajectory in Fig. 1.3 was no fluke: *every* trajectory will decay in the long run, with a growth rate approaching -5.3% per round.

There are two average wealth multipliers – $\langle r \rangle$ and r_{time} – and corresponding growth rates – $g_{\langle} \rangle$ and g_{time} – which we have determined numerically and analytically. Neither average is “wrong” *per se*. Instead, each average corresponds to a different property of the model and, therefore, answers a different question about the model. A statement like “wealth goes up, on average” is meaningless on its own. It should be countered with the question: “on what type of average?”

Finite time and finite ensembles

So far we have analysed how wealth evolution differs in the many-systems and long-time limits of the game. An observable that neatly captures these two aspects of multiplicative growth is the exponential growth rate, $g(\langle x(t) \rangle_N, \Delta t)$, observed over finite time, Δt , in a finite ensemble of N players. Here we make the growth rate an explicit function of the growing quantity and the period over which the growth is measured, because it is a random variable whose distribution depends on both. In this notation, the evolution equation for the finite-ensemble average wealth is

$$\langle x(t + \Delta t) \rangle_N = \langle x(t) \rangle_N \exp [g(\langle x(t) \rangle_N, \Delta t) \Delta t]. \quad (2.33)$$

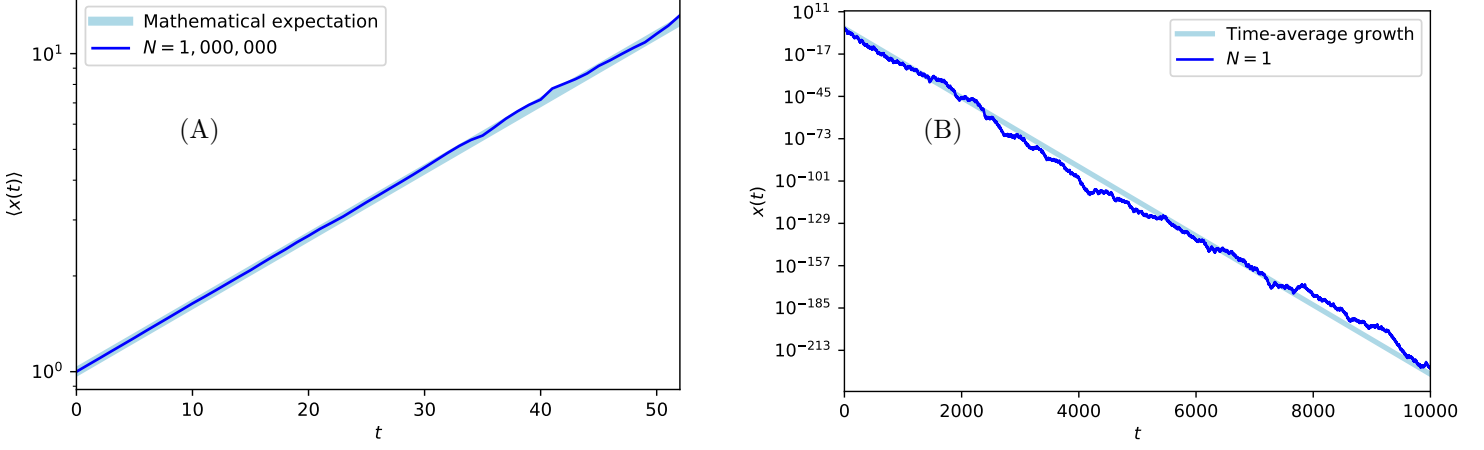


Figure 2.4: (A) Finite-ensemble average for $N = 10^6$ and 52 time steps, the light blue line is the expectation value. (B) A single system simulated for 10,000 time steps, the light blue line decays exponentially with factor r_{time} in each time step. {fig:ens_v_time}

Rearranging for the growth rate, we get

$$g(\langle x(t) \rangle_N, \Delta t) = \frac{\Delta \ln \langle x \rangle_N}{\Delta t}, \quad (2.34) \quad \{\text{eq:gest}\}$$

where the Δ in the numerator corresponds to the change over the period, Δt , in the denominator. For N and Δt finite, this is a random variable. Scalar quantities are obtained in two different limits of the same stochastic object. The exponential growth rate of the expectation value is

$$g_{\langle \rangle} = \lim_{N \rightarrow \infty} g(\langle x(t) \rangle_N, \Delta t), \quad (2.35)$$

which is also $\ln \langle r \rangle / \delta t$. The exponential growth rate followed by any single trajectory ($N = 1$) when observed for a long time is

$$g_{\text{time}} = \lim_{\Delta t \rightarrow \infty} g(x(t), \Delta t), \quad (2.36) \quad \{\text{eq:gt}\}$$

which is also $\ln r_{\text{time}} / \delta t$.

We can also write (Eq. 2.34) as a sum of the logarithmic differences in the T individual rounds of the gamble that make up the time interval $\Delta t = T\delta t$

$$g(\langle x(t) \rangle_N, \Delta t) = \frac{1}{T\delta t} \sum_{\tau=1}^T \Delta \ln \langle x(t + \tau\delta t) \rangle_N. \quad (2.37)$$

For a single trajectory, $N = 1$, in the $\Delta t \rightarrow \infty$ limit, we recover the remarkable result in (Eq. 2.31). It can be shown [49] that the time-average growth rate of a single trajectory is the same as that of a finite-ensemble average of trajectories, *i.e.*

$$\lim_{\Delta t \rightarrow \infty} \frac{\Delta \ln x}{\Delta t} = \lim_{\Delta t \rightarrow \infty} \frac{\Delta \ln \langle x \rangle_N}{\Delta t}. \quad (2.38)$$

In Sec. 6.2, we will derive this result as well as growth rates in finite ensembles and finite time.

Excursion: Scalars

Coming soon.

History: William Allen Whitworth

Coming soon.

Excursion: Compounding growth, exponentials, and the logarithm

Coming soon.

Excursion: Dimensional analysis

Coming soon.

2.7 Ergodicity

We have encountered two types of averaging – the ensemble average and the time average. In our case – assessing whether it will be good for you to play our game, the time average is the interesting quantity because it tells you what happens to your wealth as time passes. The ensemble average is irrelevant because you do not live your life as an ensemble of many yous who can average over their wealths. Whether you like it or not, you will experience yourself owning your own wealth at future times; whether you like it not, you will never experience yourself owning the wealth of a different realization of yourself. The different realizations, and therefore the expectation value, are fiction, fantasy, imagined.

We are fully aware that it can be counter-intuitive that with probability one, a different rate is observed for the expectation value than for any trajectory over time. It sounds strange that the expectation value is completely irrelevant to the problem. A reason for the intuitive discomfort is history: since the 1650s we have been trained to compute expectation values, with the implicit belief that they will reflect what happens over time. It may be helpful to point out that all of this trouble has a name that's well-known to certain people, and that an entire field of mathematics is devoted to dealing with precisely this problem. The field of mathematics is called “ergodic theory.” It emerged from the question under what circumstances the expectation value is informative of what happens over time, first raised in the development of statistical mechanics by Maxwell and Boltzmann starting in the 1850s. These lecture notes are our attempt to use precisely the insights of these physicists to re-develop economic theory from the foundations up.

History: Randomness and ergodicity in physics

Coming soon.

To convey concisely that we cannot use the expectation value and the time average interchangeably in our game, we would say “the observable x is not ergodic.”

DEFINITION: Ergodic property

In these notes, an observable A is called ergodic if its expectation value is constant in time and its time average converges to this value with probability one^a

$$\lim_{\Delta t \rightarrow \infty} \frac{1}{\Delta t} \int_t^{t+\Delta t} A(s) ds = \lim_{N \rightarrow \infty} \frac{1}{N} \sum_i^N A_i(t). \quad (2.39) \quad \{\text{eq: def_ergodic}\}$$

^aSome researchers would call A “mean ergodic” and require further observables derived from it to be (mean) ergodic in order to call A “wide-sense ergodic.” This extra nomenclature is not necessary for our work, but we leave a footnote here to avoid confusion.

The **right-hand side (RHS)** of (Eq. 2.39) is evaluated at time t , and unlike the **left-hand side (LHS)** could be a function of time. For now, we restrict our definition of ergodicity to a setup where that is not the case, *i.e.* where the ergodic property holds at all times. In Sec. 7.2.6 we will discuss transient behavior, where the distribution of A is time dependent. We then also consider an observable “ergodic” if its expectation value only converges to the time-average in the $t \rightarrow \infty$ limit.

In terms of random variables, Z , and stochastic processes, $Y_Z(t)$, the ergodic property can be visualized as in Fig. 2.2. Averaging a stochastic process over time or over the ensemble are completely different operations, and only under very rare circumstances (namely under ergodicity) can the two operations be interchanged. In our coin-tossing game the operations are clearly not interchangeable. An implicit assumption of interchangeability in the early days is the Original Sin of economic theory.

We stress that in a given setup, some observables may have the ergodic property even if others do not. Language therefore must be used carefully. Saying our game is non-ergodic really means that some key observables of interest, most notably wealth x , are not ergodic. Wealth $x(t)$, defined by (Eq. 1.1), is clearly not ergodic – with $A = x$ the LHS of (Eq. 2.39) is zero, and the RHS is not constant in time but grows. The expectation value $\langle x \rangle(t)$ simply doesn’t give us the relevant information about the temporal behavior of $x(t)$.

This does not mean that no ergodic observables exist that are related to x . Such observables do exist, and we have already encountered two of them. In fact, we will encounter a particular type of them frequently – in our quest for an observable that tells us what happens over time in a stochastic system we will find them automatically. However, again, the issue is subtle: an ergodic observable may or may not tell us what we’re interested in. It may be ergodic but not indicate what happens to x . For example, the multiplicative factor $r(t)$ is an ergodic observable that reflects what happens to the expectation value of x , whereas per-round changes in the logarithm of wealth, $\delta \ln x = \ln r$,

are also ergodic and reflect what happens to x over time. **This reflects the importance of the observable appearing meaningfully in a sum.**

Proposition: $r(t)$ and $\delta \ln x$ are ergodic for the wealth dynamic defined by (Eq. 1.1) and (Eq. 1.2).

Proof. According to (Eq. 2.10) and (Eq. 2.9), the expectation value of $r(t)$ is

$$\langle r \rangle = \lim_{N \rightarrow \infty} \frac{1}{N} \sum_i^N r_i, \quad (2.40) \quad \{\text{eq:e_r}\}$$

and, according to (Eq. 2.16), the time average of $r(t)$ is

$$\bar{r} = \lim_{T \rightarrow \infty} \frac{1}{T} \sum_{\tau}^T r_{\tau}, \quad (2.41) \quad \{\text{eq:t_r}\}$$

where we have written $r_{\tau} = r(t + \tau \delta t)$ to make clear the equivalence between the two expressions. The only difference is between the labels we have chosen for the dummy variable (i in (Eq. 2.40) and τ in (Eq. 2.41)). Clearly, the expressions yield the same value.

The same argument holds for $\delta \ln x$. □

Whether we consider (Eq. 2.41) an average over time or over an ensemble is only a matter of our interpretation. The mathematics are the same.

The expectation value $\langle \delta \ln x \rangle$ is important, historically. Daniel Bernoulli noticed in 1738 [6] that people tend to optimize $\langle \delta \ln x \rangle$, whereas it had been assumed that they should optimize $\langle \delta x \rangle$. Unaware of the issue of ergodicity (200 years before the concept was discovered and the word was coined), Bernoulli had no good explanation for this empirical fact and simply stated that people tend to behave as though they valued money non-linearly. We now know what is actually going on: multiplicative dynamics are a fairly realistic model for real wealth, and under those dynamics δx is not ergodic, and $\langle \delta x \rangle$ is of no interest – it doesn't tell us what happens over time. However, $\delta \ln x$ *is* ergodic, and $\langle \delta \ln x \rangle$ does tell us what happens to x over time, wherefore seeing people optimise $\langle \delta \ln x \rangle$ just means seeing them optimise wealth over the one trajectory that describes a financial life, rather than across the ensemble of possibilities.

Ergodicity is not the same concept as stationarity. As an illustration of the difference, consider the following process: $f(t) = z_i$, where z_i is an instance of a random variable Z . Explicitly, this means a realisation of the stochastic process $f(t)$ is generated as follows: we generate the random instance z_i once, and then fix $f(t)$ at that value for all time. The distribution of $f(t)$ is independent of t and in that sense $f(t)$ is stationary. But it is not ergodic: averaging over the ensemble, we obtain $\langle f(t) \rangle = \langle z \rangle$, whereas averaging over time in the i^{th} trajectory gives $\bar{f} = z_i$. Thus the process is stationary but not ergodic.

2.8 Changes and stability

Coming soon. Most material now appears in the next chapter.

`{section:Rates}`

2.9 Normal distribution

{section:Normal_distribut

Coming soon.

2.10 Brownian motion

{section:Brownian_motion}

We motivate the model called **Brownian motion (BM)** as a limiting process, the continuous-time limit, that arises from random walks. In the previous section we established that the discrete increments of the logarithm of x , which we called v , are instances of a time-independent random variable in our game. A quantity making such random steps over time is said to perform a “random walk.” Indeed, the blue line for a single system in Fig. 1.2 (B) shows 52 steps of a random walk trajectory. Random walks come in many forms – in all of them v changes discontinuously by an amount δv drawn from a time-independent distribution, over time intervals which may be regular or which may be drawn from a time-independent distribution themselves.

We are interested only in the simple case where v changes at regular intervals, $\delta t, 2\delta t, \dots$. For the distribution of increments we only insist on the existence of the variance, meaning we insist that $\text{var}(\delta v) = \langle \delta v^2 \rangle - \langle \delta v \rangle^2$ be finite. Increments whose distributions are heavier-tailed do not lead to BM (BM has continuous paths, and that continuity is broken by such increments).

The change in v after a long time is the sum of many independent increments,

$$v(t + T\delta t) - v(t) = \sum_i^T \delta v_i. \quad (2.42)$$

The Gaussian central limit theorem tells us that such a sum will become Gaussian-distributed as we add more terms to the sum and re-scale it appropriately, namely so as to keep the width of the distribution finite and remove any systematic drift,

$$\lim_{T \rightarrow \infty} \frac{1}{\sqrt{T}} \sum_i^T (\delta v_i - \underbrace{\langle \delta v \rangle}_{\text{remove systematic drift}}) \sim \mathcal{N}(0, \text{var}(\delta v)), \quad (2.43) \quad \{\text{eq:CLT}\}$$

where we call $\frac{\langle \delta v \rangle}{\delta t}$ the “drift term.” The notation $\sim \mathcal{N}(0, \text{var}(\delta v))$ is short-hand for “is Gaussian distributed, with mean 0 and variance $\text{var}(\delta v)$.” The logarithmic change in the long-time limit that was of interest to us in the analysis of the coin toss game is thus Gaussian distributed.

Let’s also ask about the re-scaling that was applied in (Eq. 2.43). Scaling properties are very robust, and especially the scaling of random walks for long times will be useful to us.

We work with the simplest setup: at time zero we start at zero, $v(0) = 0$, and in each time step, we either increase or decrease v by 1, with probability 1/2. To avoid notation clutter, we’ll set the duration of a time step to $\delta t = 1$, so that T is both the number of steps and the time.

We are interested in the variance of the distribution of v as T increases, which we obtain by computing the first and second moments of the distribution.

The first moment (the expectation value) of v is $\langle v \rangle(T) = 0$, by symmetry for all times.

We obtain the second moment by induction²: Whatever the second moment, $\langle v(T)^2 \rangle$, is at time T , we can write down its value at time $T + 1$ as

$$\langle v(T+1)^2 \rangle = \frac{1}{2} [\langle (v(T)+1)^2 \rangle + \langle (v(T)-1)^2 \rangle] \quad (2.44)$$

$$= \frac{1}{2} [\langle v(T)^2 + 1 + 2v(T) \rangle + \langle v(T)^2 + 1 - 2v(T) \rangle] \quad (2.45)$$

$$= \langle v(T)^2 \rangle + 1. \quad (2.46)$$

In addition, we know the initial value of $v(0) = 0$. By induction it follows that the second moment is

$$\langle v(T)^2 \rangle = T \quad (2.47)$$

and, since the first moment is zero, the variance is

$$\text{var}(v(T)) = T. \quad (2.48) \quad \{\text{eq:BM_var}\}$$

The standard deviation – the width of the distribution – of changes in a quantity following a random walk thus scales as the square-root of the number of steps that have been taken, $\sqrt{T}$.

This square-root behaviour leads to many interesting results. It can make averages stable (because $\sqrt{T}/T$ converges to zero for large T), and sums unstable (because $\sqrt{T}$ diverges for large T). Consequently, we may expect that as the size of some system increases, some properties become stable and others unstable.

Imagine simulating a single long trajectory of v and plotting it on paper³. The amount of time that has to be represented by a fixed length of paper increases linearly with the simulated time because the paper has a finite width to accommodate the horizontal axis. If $\langle \delta v \rangle \neq 0$ then the amount of variation in v that has to be represented by a fixed amount of paper also increases linearly with the simulated time. However, the departures of Δv from its expectation value $T \langle \delta v \rangle$ only increase as the square-root of T . Thus, the amount of paper-space given to these departures scales as $T^{-1/2}$, and for very long simulated times the trajectory will look like a straight line on paper.

In an intermediate regime, fluctuations will still be visible but they will also be approximately Gaussian distributed. In this regime it is often easier to replace the random walk model with the corresponding continuous process. That process – finally – is BM.

We think of BM as the limit of a random walk where we shorten the duration of a step $\delta t \rightarrow 0$, and scale the width of an individual step so as to maintain the random-walk scaling of the variance, meaning $|\delta v| = \sqrt{\delta t}$. In the limit $\delta t \rightarrow 0$, this implies that the local slope of a BM trajectory diverges, $\frac{\delta v}{\delta t} \rightarrow \infty$. This means that BM trajectories are infinitely jagged, or – in mathematical terms – they are not differentiable. However, the way in which they become non-differentiable, through the $\sqrt{\delta t}$ factor, just leaves the trajectories continuous (this isn't the case for $|\delta v| = \delta t^\alpha$, where α is less than 0.5).

²The argument is nicely illustrated in [19, Volume 1, Chapter 6-4], where we first came across it.

³This argument is inspired by a colloquium presented by Wendelin Werner in the mathematics department of Imperial College London in January 2012. Werner started the colloquium with a slide that showed a straight horizontal line and asked: what is this? Then answered that it was the trajectory of a random walk, with the vertical and horizontal axes scaled equally.

Continuity of v means that it is possible to make the difference $|v(t) - v(t + \epsilon)|$ arbitrarily small by choosing ϵ sufficiently small. Trajectories (of non-BM processes) that don't have this property contain what are appropriately called "jumps." Continuity therefore means that there are no jumps. These subtleties make BM a topic of great mathematical interest, and many books have been written about it. We will pick from these books only what is immediately useful to us. To convey the universality of BM we define it formally as follows:

DEFINITION: Brownian motion i

If a stochastic process has continuous paths, stationary independent increments, and is distributed according to $\mathcal{N}(\mu t, \sigma^2 t)$ then it is a Brownian motion.

The process can be defined in different ways. Another illuminating definition is this:

DEFINITION: Brownian motion ii

If a stochastic process is continuous, with stationary independent increments, then the process is a Brownian motion.

We quote from [20]: *"This beautiful theorem shows that Brownian motion can actually be defined by stationary independent increments and path continuity alone, with normality following as a consequence of these assumptions. This may do more than any other characterization to explain the significance of Brownian motion for probabilistic modeling."*

Indeed, BM is not just a mathematically rich model but also – due to its emergence through the Gaussian central limit theorem – a model that represents a large universality class, *i.e.* it is a good description of what happens over long times in many other models that produce random trajectories.

The power of BM lies in its simplicity and analytic tractability, involving only two parameters, μ and σ . We will often work with its representation as a stochastic differential equation (SDE)

$$dv = \mu dt + \sigma dW \quad (2.49) \quad \{\text{eq:BM_dx}\}$$

where dW is the so-called "Wiener increment," the beating heart of many SDEs. The Wiener increment can be defined by two properties: its distribution and its auto-correlation,

$$dW \sim \mathcal{N}(0, dt) \quad (2.50)$$

$$\langle dW(t)dW(t') \rangle = dt \delta(t, t'), \quad (2.51)$$

where $\delta(t, t')$ is the Kronecker delta – zero if its two arguments differ ($t \neq t'$), and one if they are identical ($t = t'$).⁴ In simulations BM paths can be constructed from a discretized version of (Eq. 2.49)

$$v(t + \delta t) = v(t) + \mu \delta t + \sigma \sqrt{\delta t} \xi_t, \quad (2.52) \quad \{\text{eq:BM_d}\}$$

⁴Physicists often write $dW = \eta dt$, where $\langle \eta \rangle = 0$ and $\langle \eta(t)\eta(t') \rangle = \delta(t - t')$, in which case $\delta(t - t')$ is the Dirac delta function, defined by the integral $\int_{-\infty}^{\infty} f(t)\delta(t - t')dt = f(t')$. Because of its singular nature ($\eta(t)$ does not exist ("is infinite"), only its integral exists) it can be difficult to develop an intuition for this object, and we prefer the dW notation.

where ξ_t are instances of a standard normal distribution ($\mathcal{N}(0, 1)$).

BM itself is not a time-independent random variable – it is a non-ergodic stochastic process. This is easily seen by comparing expectation value and time average. We start with the expressions (stated without proof here) for the finite-ensemble average and the finite-time average of BM. The finite-ensemble average (easy to derive) is distributed as

$$\langle v \rangle_N \sim \mu t + \mathcal{N}(0, t/N), \quad (2.53) \quad \{\text{eq:fin_ens_BM}\}$$

and the finite-time average (a little harder to derive) of a single BM trajectory is distributed as

$$\bar{v}_t \sim \mu t/2 + \sigma \mathcal{N}(0, t/3). \quad (2.54) \quad \{\text{eq:fin_tim_BM}\}$$

The expectation value, *i.e.* the limit $N \rightarrow \infty$ of (Eq. 2.53), converges to μt with probability one, so it depends on time, and it's unclear how to compare that to a time average (which cannot depend on time). Its limit $t \rightarrow \infty$ does not exist.

The time average, the limit $t \rightarrow \infty$ of (Eq. 2.54) diverges unless $\mu = 0$, but even with $\mu = 0$ the limit is a random variable with diverging variance – something whose density is zero everywhere. In no meaningful sense do the two expressions converge to the same scalar in the relevant limits.

Clearly, BM, whose increments are ergodic, is itself not ergodic. However, that doesn't make it unmanageable or unpredictable – we know the distribution of BM at any moment in time. But the non-ergodicity has surprising consequences of which we mention one now. We already mentioned that if we plot a Brownian trajectory with non-zero drift on a piece of paper it will turn into a straight line for long enough simulation times. This suggests that the randomness of a Brownian trajectory becomes irrelevant under a very natural rescaling. Inspired by this insight let's hazard a guess as to what the time-average of zero-drift BM might be.

The simplest form of zero-drift BM starts at zero, $v(0) = 0$ and has variance $\text{var}(v(t)) = t$ (this process is also known as the “Wiener process”). The process is known to be recurrent – it returns to zero, arbitrarily many times, with probability one in the long-time limit. We would not be mad to guess that the time average of zero-drift BM,

$$\bar{v} = \lim_{t \rightarrow \infty} \frac{1}{t} \int_0^t dt' v(t'), \quad (2.55)$$

will converge to zero with probability one. But we would be wrong. Yes, the process has no drift, and yes it returns to zero infinitely many times, but its time average is not a delta function at zero. It is, instead normally distributed with infinite variance according to the following limit

$$\bar{v} \sim \lim_{t \rightarrow \infty} \mathcal{N}(0, t/3). \quad (2.56)$$

Averaging over time, in this case, does not remove the randomness. A sample trajectory of the finite-time average (not of BM but of the average over a BM) is shown in Fig. 2.5. In the literature this process, $\frac{1}{t} \int_0^t dt' v(t')$, is known as the “random acceleration process” [11].

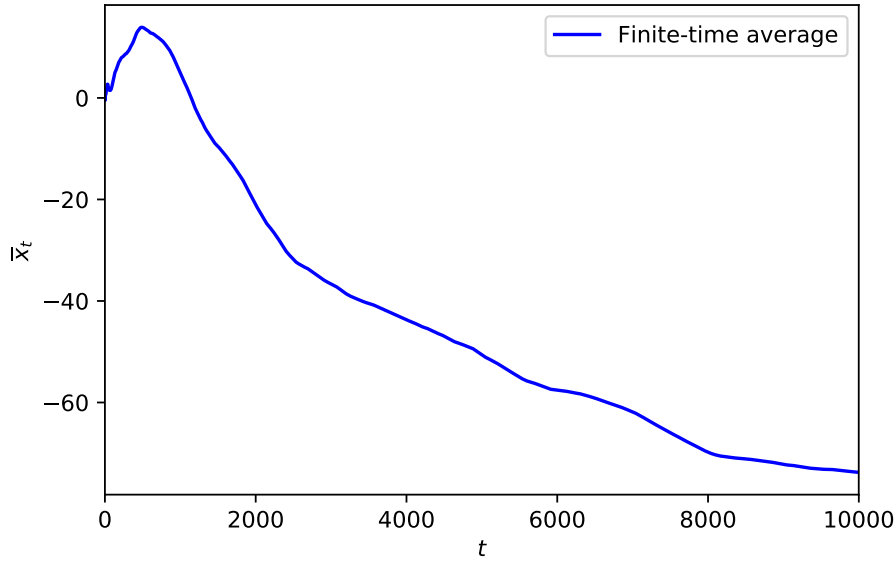


Figure 2.5: Trajectory of the finite-time average of a zero-drift BM. The process is not ergodic: the time average does not converge to a number, but is instead distributed according to $\mathcal{N}(0, t/3)$ for all times, while the expectation value is zero. It is the result of integrating a BM; integration is a smoothing operation, and as a consequence the trajectories are smoother than BM (unlike a BM trajectory, they are differentiable).

{fig:1_6}

2.11 Geometric Brownian motion

{section:Geometric_Browni}

DEFINITION: Geometric Brownian motion

If the logarithm of a quantity performs Brownian motion, the quantity itself performs “geometric Brownian motion.”

While in Sec. 2.10 $v(x) = \ln(x)$ performed BM, x itself performed geometric Brownian motion (GBM). The change of variable from x to $v(x) = \ln(x)$ is trivial in a sense but it has interesting consequences. It implies, for instance, that

- $x(t)$ is log-normally distributed
- increments in x are neither stationary nor independent
- $x(t)$ cannot become negative
- the most likely value of x (the mode) does not coincide with the expectation value of x .

These and other properties of the log-normal distribution will be discussed in detail in Sec. 6.1.1.

Again, it is informative to write GBM as a stochastic differential equation.

$$dx = x(\mu dt + \sigma dW). \quad (2.57) \quad \{\text{eq:GBM_c}\}$$

Similarly to [BM](#), trajectories for [GBM](#) can be simulated using the discretized form (*cf.* (Eq. 2.52))

$$\delta x = x(\mu\delta t + \sigma\sqrt{\delta t}\xi_t), \quad (2.58) \quad \{\text{eq:GBM_d}\}$$

where $\xi_t \sim \mathcal{N}(0, 1)$ are instances of a standard normal variable. In such simulations we must pay attention that the discretization does not lead to negative values of x . This happens if the expression in brackets in (Eq. 2.58) is smaller than -1 (in which case x changes negatively by more than itself). To avoid negative values we must have $\mu\delta t + \sigma\sqrt{\delta t}\xi_t > -1$, or $\xi_t < \frac{1+\mu\delta t}{\sigma\sqrt{\delta t}}$. As δt becomes large it becomes more likely for ξ_t to exceed this value, in which case the simulation fails. But ξ_t is Gaussian distributed, meaning it has thin tails, and choosing a sufficiently small value of δt makes these failures essentially impossible.

[GBM](#) on logarithmic vertical scales looks like [BM](#) on linear vertical scales. Figure 1.2 is, in fact, an example of a very coarse discretisation of [GBM](#). But it's useful to look at a more finely discretised trajectory of [GBM](#) on linear scales to develop an intuition for this important process.

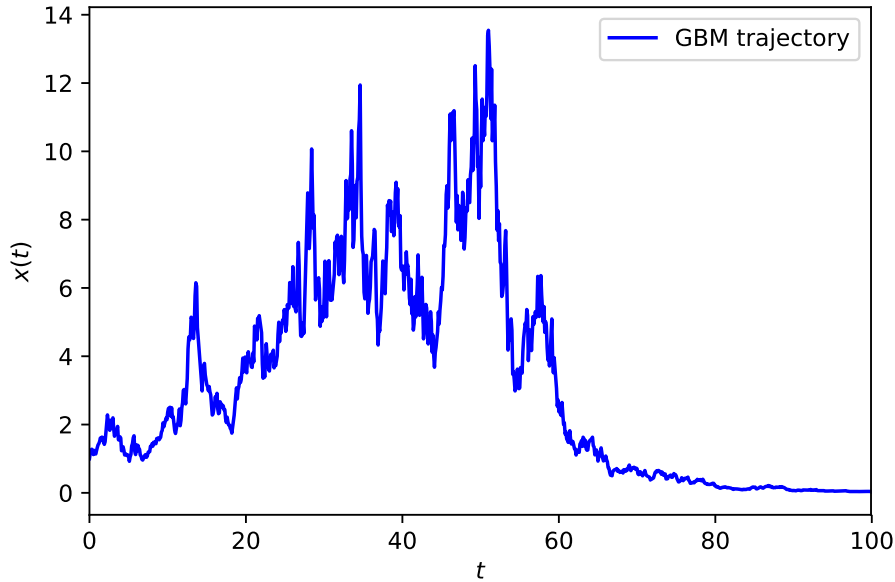


Figure 2.6: Trajectory of a [GBM](#). What happens to the trajectory tomorrow depends strongly on where it is today – for instance, unlike for [BM](#), it is difficult to recover from a low value of x , and trajectories are likely to get stuck near zero. Occasional excursions are characterised by large fluctuations. Parameters are $\mu = 0.05$ per time unit and $\sigma = \sqrt{2\mu}$, corresponding to zero growth rate in the long run. It would be easy to invent a story to go with this (completely random) trajectory – perhaps something like “things were going well in the beginning but then a massive crash occurred that destroyed morale.”

{fig:1_7}

The basic message of the game from Sec. 1.1 is that we may obtain different values for growth rates, depending on how we average – an expectation value is one average, a time average is quite another. The game itself is sometimes called the multiplicative binomial process [51], we thank S. Redner for pointing

this out to us. [GBM](#) is the continuous version of the multiplicative binomial process, and it shares the basic feature of a difference between the growth rate of the expectation value and time-average growth.

The expectation value is easily computed – the process is not ergodic, but that does not mean we cannot compute its expectation value. We simply take the expectations values of both sides of (Eq. 2.57) to get

$$\langle dx \rangle = \langle x(\mu dt + \sigma dW) \rangle \quad (2.59)$$

$$= d \langle x \rangle = \langle x \rangle \mu dt. \quad (2.60)$$

This differential equation has the solution

$$\langle x(t) \rangle = x(t_0) \exp(\mu t), \quad (2.61)$$

which determines the growth rate of the expectation value as

$$g_m(\langle x \rangle) = \mu. \quad (2.62) \quad \{\text{eq:expectation_g}\}$$

As we know, this growth rate is different from the growth rate that materializes with probability 1 in the long run. Computing the time-average growth rate is only slightly more complicated, and it will get even simpler once we've introduced Itô calculus in Sec. 2.12. But for now we will follow this plan: consider the discrete process (Eq. 2.58) and compute the changes in the logarithm of x , then we will let δt become infinitesimal and arrive at the result for the continuous process. We know $\delta \ln(x(t))$ to be ergodic and reflective of performance over time, wherefore we will proceed to take its expectation value to compute the time average of the exponential growth rate of the process.

The change in the logarithm of x in a time interval δt is

$$\ln x(t + \delta t) - \ln x(t) = \ln[x(1 + \mu\delta t + \sigma\sqrt{\delta t}\xi_t)] - \ln x(t) \quad (2.63)$$

$$= \ln x + \ln(1 + \mu\delta t + \sigma\sqrt{\delta t}\xi_t) - \ln x(t) \quad (2.64)$$

$$= \ln(1 + \mu\delta t + \sigma\sqrt{\delta t}\xi_t), \quad (2.65)$$

which we Taylor-expand as $\ln(1 + \text{something small})$ because we will let δt become small. Expanding to second order,

$$\ln x(t + \delta t) - \ln x(t) = \mu\delta t + \sigma\sqrt{\delta t}\xi_t - \frac{1}{2} \left(\mu\sigma\delta t^{3/2}\xi_t + \sigma^2\delta t\xi_t^2 \right) + o(\delta t^2), \quad (2.66)$$

using “little-o notation” to denote terms that are of order δt^2 or smaller. Finally, because $\delta \ln x(t)$ is ergodic, by taking the expectation value of this equation we find the time average of $\delta \ln x(t)$

$$\langle \ln x(t + \delta t) - \ln x(t) \rangle = \mu\delta t - \frac{1}{2} (\mu^2\delta t^2 + \sigma^2\delta t) + o(\delta t^2). \quad (2.67)$$

Letting δt become infinitesimal the higher-order terms in δt vanish, and we find

$$\langle \ln x(t + dt) - \ln x(t) \rangle = \mu dt - \frac{1}{2} \sigma^2 dt, \quad (2.68)$$

so that the time-average growth rate is

$$\bar{g} = \frac{d \langle \ln x \rangle}{dt} = \mu - \frac{1}{2} \sigma^2. \quad (2.69) \quad \{\text{eq:time_g}\}$$

The non-ergodicity of GBM leads to a difference between the behaviour of the expectation value (which grows at $g_m(\langle x \rangle)$) and the long-time behaviour of any given trajectory (which grows at $\bar{g}$). Because people experience their wealth over time (which may be described by GBM) and have not access to the ensemble of other possible trajectories, they quite reasonably behave closer to optimising $g_m(\langle x \rangle)$ than to $\bar{g}$.

We could have guessed the result by combining Whitworth's argument on the disadvantage of gambling with the scaling of BM. Let's re-write the factor $1 - \epsilon$ in (Eq. ??) as $1 - \sigma\sqrt{\delta t}$. According to the scaling of the variance in a random walk, (Eq. 2.48), this would be a good coarse-graining of some faster process (with shorter time step) underlying Whitworth's game. To find out what happens over one single time step we take the square root of (Eq. ??),

$$[(1 + \sigma\sqrt{\delta t})(1 - \sigma\sqrt{\delta t})]^{1/2} = [1 - \sigma^2\delta t]^{1/2}. \quad (2.70)$$

Letting δt become infinitesimally small, we replace δt by dt , and the first-order term of a Taylor-expansion becomes exact,

$$[(1 + \sigma\sqrt{\delta t})(1 - \sigma\sqrt{\delta t})]^{1/2} \rightarrow 1 - \frac{\sigma^2}{2}dt, \quad (2.71)$$

in agreement with (Eq. 2.69) if the drift term $\mu = 0$, as assumed by Whitworth.

2.12 Itô calculus

{section:Itô}

We have chosen to work with the discrete process here and have arrived at a result that is more commonly shown using Itô's formula. We will not discuss Itô calculus in depth but we will use some of its results. The key insight of Itô was that the non-differentiability of so-called Itô processes leads to a new form of calculus, where in particular the chain rule of ordinary calculus is replaced. An Itô process is a SDE of the following form

$$dx = a(x, t)dt + b(x, t)dW. \quad (2.72) \quad \{\text{eq:Itô_process}\}$$

If we are interested in the behaviour of some other quantity that is a function of x , let's say $v(x)$, then Itô's formula tells us how to derive the relevant SDE as follows:

$$dv = \left(\frac{\partial v}{\partial t} + a(x, t)\frac{\partial v}{\partial x} + \frac{b(x, t)^2}{2}\frac{\partial^2 v}{\partial x^2} \right) dt + b(x, t)\frac{\partial v}{\partial x}dW. \quad (2.73) \quad \{\text{eq:Itô}\}$$

Derivations of this formula can be found on Wikipedia. Intuitive derivations, such as [21], use the scaling of the variance, (Eq. 2.48), and more formal derivations, along the lines of [20], rely on integrals. We simply accept (Eq. 2.73) as given. It makes it very easy to re-derive (Eq. 2.69), which we leave as an exercise: use (Eq. 2.73) to find the SDE for $\ln(x)$, take its expectation value and differentiate with respect to t . We will use (Eq. 2.73) in Sec. 6.1.5. The above computations are intended to give the reader intuitive confidence that Itô calculus can be trusted⁵. We find that, though phrased in different words, our

⁵Itô calculus is one way of interpreting the non-differentiability of dW . Another interpretation is due to Stratonovich, which is not strictly equivalent. However, the key property of GBM that we make extensive use of is the difference between the growth rate of the expectation value, $g_m(\langle x \rangle)$, and the time-average growth rate, $\bar{g}$. This difference is the same in the Stratonovich and the Itô interpretation, and all our results hold in both cases.

key insight – that *the growth rate of the expectation value is not the time-average growth rate* – has appeared in the literature not only in 1870 but also in 1944. And in 1956 [24], and in 1966 [59], and in 1991 [14], and at many other times. Yet the depth of this insight remained unprobed.

Equation (2.69), which agrees with Itô calculus, may be surprising. Consider the case of no noise $dx = x\mu dt$. Here we can identify $\mu = \frac{1}{x} \frac{dx}{dt}$ as the infinitesimal increment in the logarithm, $\frac{d \ln(x)}{dt}$, using the chain rule of ordinary calculus. A naïve application of the chain rule to (Eq. 2.57) would therefore also yield $\frac{d \langle \ln(x) \rangle}{dx} = \mu$, but the fluctuations in GBM have a non-linear effect, and it turns out that the usual chain rule does not apply. Itô calculus is a modified chain rule, (Eq. 2.73), which leads to the difference $-\frac{\sigma^2}{2}$ between the expectation-value growth rate and the time-average growth rate.

This difference is sometimes called the “spurious drift”, but at the [London Mathematical Laboratory \(LML\)](#) we call it the “Weltschmerz” because it is the difference between the many worlds of our dreams and fantasies, and the one cruel reality that the passage of time imposes on us.

Summary of Chap. 2

Coming soon.

Part II

Microeconomics

Chapter 3

Decisions in a riskless world

{chapter:Riskless}

Decision theory is a cornerstone of formal economics. As the name suggests, it models how people make decisions. In this chapter we will generalise and formalise the treatment of the coin tossing game to introduce our approach to decision theory. Our central axiom will be that people attempt to maximize the rate at which wealth grows when averaged over time. This is a surprisingly powerful idea. In many cases it eliminates the need for well established but epistemologically troublesome techniques, such as utility functions.

3.1 Models and science fiction

{section:Models_and}

We will do decision theory by using mathematical models, and since this can be done in many ways we will be explicit about how we choose to do it. We will define a wealth process – a model of how wealth changes with time – and a decision criterion. The wealth process and the decision criterion may or may not remind you of the real world. We will not worry too much about the accuracy of these reminiscences. Instead we will “shut up and calculate” – we will let the mathematical model create its world. Writing down a mathematical model is like laying out the premise for a science-fiction novel. We may decide that people can download their consciousness onto a computer, that medicine has advanced to eliminate ageing and death – these are premises we are at liberty to invent. Once we have written them down we begin to explore the world that results from those premises. A decision criterion is really a model of human behaviour – what makes us who we are if not our decisions? It therefore implies a long list of specific behaviours that will be observed in a given model world. For example, some criteria will lead to cooperation, others will not, some will lead to the existence of insurance contracts, others will not *etc.* We will explore the worlds created by the different models. Once we have done so we invite you to judge which model you find most useful for your understanding of the world. Of course, having spent many years thinking about these issues we have come to our own conclusions, and we will put them forward because we believe them to be helpful.

To keep the discussion to a manageable volume we will only consider a setup that corresponds to making purely financial decisions. We may bet on a horse or take out personal liability insurance. This chapter will not tell you whom you should marry or even whose economics lectures you should attend.

3.2 The decision axiom

A “decision theory” is a model of human behaviour. We will write down such a model phrased as the following simple axiom:

Decision axiom

People optimize the growth rate of their wealth.

Without discussing why people might do this, let’s step into the world created by this axiom. To do that, we need to be crystal clear about what a growth rate is, so we’ll discuss that first, in Sec. 3.3. Traditionally, decision theory deals with an uncertain future: we have to decide on a course of action now although we don’t know with certainty what will happen to us in the future under any of our choices. We will systematically work our way towards this setup, beginning with trivial decisions where neither time nor uncertainty matters Sec. 3.4.1, next introducing time Sec. 3.4.2 (where we will shed light on what’s called “discounting”). In the next chapter we will introduce uncertainty, see Sec. ?? (where we will shed light on what’s called “expected utility theory”).

3.3 Growth rates

{section:Growth_rates}

You may have wondered why both

$$g_a = \frac{x(t + \Delta t) - x(t)}{\Delta t} \quad (3.1) \quad \{\text{eq:add_rate}\}$$

and

$$g_e = \frac{\ln x(t + \Delta t) - \ln x(t)}{\Delta t} \quad (3.2) \quad \{\text{eq:exp_rate}\}$$

are sometimes called a growth rate – they’re different objects, why the same name? By the end of this section, the answer to this question should be clear.

When we say that $x(t)$ is a growth process, we mean that it is a monotonic function of t . If you’re thinking about randomness – don’t, we’ll come to that later. For now, we will just work with a deterministic function $x(t)$ – we even use an unusual font to indicate that this is a deterministic function.

For a given process, the appropriate growth rate, g , solves the following problem for us: how do we characterise how fast x grows? A growth rate is a mathematical object of the form

$$g = \frac{\Delta v(x)}{\Delta t}, \quad (3.3) \quad \{\text{eq:gen_rate}\}$$

where $v(x)$ is a monotonically increasing function of wealth x . Comparing to (Eq. 3.1) and (Eq. 3.2), we find that $v(x) = x$ for additive dynamics and $v(x) = \ln x$ for multiplicative dynamics. The transformation $x \rightarrow v(x)$ ensures temporal stability of g .

How does this work, and what does it mean? Specifically,

1. why the transformation?
2. how do we know which transformation to use?

We will start with the mathematically simple case of the additive growth rate, (Eq. 3.1), and discuss its properties by applying it to additive growth. Next we will ask under what conditions the exponential growth rate, (Eq. 3.2), is appropriate, and that will lead us to the general growth rate, (Eq. 3.3).

3.3.1 Additive growth rate

{section:add_rate}

If I want to know how fast $x(t)$ grows, the most obvious thing to compute is its rate of change – that’s the additive growth rate, (Eq. 3.1). This tells me by how much x grows in the interval $[t, t + \Delta t]$. If $x(t)$ is linear in t , so that

$$x(t) = x(0) + \gamma t \quad (3.4) \quad \{\text{eq:linx}\}$$

then this is a very informative quantity. We’ll now state carefully why it is informative in this case. That may seem pedantic at this point, but it will become useful when we generalise in Sec. 3.3.2 and Sec. 3.3.3.

The additive growth rate (Eq. 3.1) is informative of how fast x grows under additive dynamics (Eq. 3.4) because in this case the t -dependence drops out: we can measure g_a whenever we want, and we’ll always get the same value, $g_a = \gamma$. Not to get too philosophical about it, but this kind of time-translation invariance

(fancy word) is a key concept in science: the search for laws is the search for universal structure – especially for time-translation invariant structure, for something “timeless.”

Let’s re-write the linear dynamic (Eq. 3.4) in differential form

$$dx = \gamma dt \quad (3.5)$$

Because γ depends neither on t nor on x , we can re-write this as

$$dx = d(\gamma t) \quad (3.6)$$

This way of writing it tells us that the growth rate γ is really a sort of clock speed. There’s no difference between rescaling t and rescaling γ (by the same factor).

We make a mental note: *the growth rate is a clock speed*. But what kind of clock speed are we talking about? What’s a clock speed anyway?

Or: what’s a clock? A clock is a process that we believe does something repeatedly at regular intervals. It lets us measure time by counting the repetitions. By convention, after 9,192,631,770 cycles of the radiation produced by the transition between two levels of the caesium 133 atom we say “one second has elapsed.” That’s just something we’ve agreed on. But any other thing that does something regularly would work as a clock – like the Earth completing one full rotation around its axis *etc.*

When we say “the growth rate of the process is γ ,” we mean that x advances by γ units on the process-scale (meaning in x) in one standard time unit (in finance we often choose one year as the unit, Earth going round the Sun). So it’s a conversion factor between the time scales of a standard clock and the process clock.

Of course, a clock is no good if it speeds up or slows down. For processes other than additive growth we have to be quite careful before we can use them as clocks, i.e. before we can state their growth rates.

3.3.2 Exponential growth rate

{section:exp_rate}

Now what about the exponential growth rate, (Eq. 3.2)? This first thing to notice is that it’s not time-translation invariant for additive growth, (Eq. 3.4). Substituting (Eq. 3.4) in (Eq. 3.2) gives

$$g_e = \frac{\ln x(t + \Delta t) - \ln x(t)}{\Delta t} \quad (3.7)$$

$$= \frac{\ln [x(t) + \gamma \Delta t] - \ln x(t)}{\Delta t} \quad (3.8)$$

$$= \frac{\ln \left(1 + \frac{\gamma \Delta t}{x(t)} \right)}{\Delta t}. \quad (3.9) \quad \{\text{eq:exp_lin}\}$$

That means the exponential growth rate does not extract the clock speed γ from linear growth. There’s a mismatch between the process and the form of the rate with which we’re measuring its speed. The exponential growth rate of additive growth is not a constant but (see RHS of (Eq. 3.9)) depends on $x(t)$, i.e. on the time when we started measuring. It also depends on how long we measured, Δt . If we used it to characterise the growth in described by (Eq. 3.4), we would

find lots of contradictions – some people would say the growth is faster, others slower because they measured at different times or for different periods.

But the exponential growth rate is commonly used, and for good reasons. Let’s see what it’s good for, by imposing that it’s useful and then working backwards to find the process we should use it for (we expect to find exponential growth).

We require that (Eq. 3.2) yield a constant, let’s call that γ again, irrespective of when we measure it.

$$g_e = \frac{\Delta \ln x}{\Delta t} = \gamma, \quad (3.10)$$

or

$$\Delta \ln x = \gamma \Delta t, \quad (3.11)$$

or indeed, in differential form, and revealing that again the growth rate is a clock speed: γ plays the same role as t ,

$$d \ln x = d(\gamma t). \quad (3.12)$$

This differential equation can be directly integrated and has the solution

$$\ln x(t) - \ln x(0) = \gamma t. \quad (3.13)$$

We solve for the dynamic $x(t)$ by writing the log difference as a fraction

$$\ln \left[\frac{x(t)}{x(0)} \right] = \gamma t, \quad (3.14)$$

and exponentiating

$$x(t) = x(0) \exp(\gamma t) \quad (3.15) \quad \{\text{eq:expx}\}$$

As expected, we find that the *exponential* growth rate, (Eq. 3.2), is the appropriate growth rate (meaning time-independent) for *exponential* growth.

In terms of clocks, what just happened is this: we insisted that (Eq. 3.2) be a good definition of a clock speed. That requires it to be constant, meaning that the process has to advance on the logarithmic scale, specified in (Eq. 3.2), by the same amount in every time interval (measured on the standard clock, of course – Earth or caesium).

3.3.3 General growth rate

{section:gen_rate}

Finally let’s be more ambitious and posit a general process, $x(t)$, of which we only assume that it grows according to a dynamic that can be written down as a separable differential equation. We could be even more general, but this is bad enough.

How do we define a growth rate now?

Well, as in Sec. 3.3.2, we insist that the thing we’re measuring will be a clock speed, *i.e.* a time-independent rescaling of time. We enforce this by writing down the dynamic in differential form, containing the growth rate as a time rescaling factor. Then we’ll work backwards and solve for g :

$$dx = f(x) d(gt) \quad (3.16) \quad \{\text{eq:gen_diff_x}\}$$

(for linear growth, like in (Eq. 3.4), $f(x)$ would just be $f(x) = 1$, and for exponential growth, (Eq. 3.15), it would be $f(x) = x$, but we're leaving it general). We separate variables in (Eq. 3.16) and integrate the differential equation

$$\int_{x(t)}^{x(t+\Delta t)} \frac{1}{f(x)} dx = g\Delta t, \quad (3.17)$$

and we've got what we want, namely the functional form of g :

$$g = \frac{\int_{x(t)}^{x(t+\Delta t)} \frac{1}{f(x)} dx}{\Delta t}. \quad (3.18) \quad \{\text{eq:g_int}\}$$

This doesn't quite look like our stated aim: the general expression for a growth rate, (Eq. 3.3). But we get there, by simplifying (Eq. 3.18) and denoting the definite integral with the letter v , so that

$$\Delta v = \int_{x(t)}^{x(t+\Delta t)} \frac{1}{f(x)} dx. \quad (3.19) \quad \{\text{eq:Dv}\}$$

Equation (3.3) now follows exactly by substituting (Eq. 3.19) in (Eq. 3.18). This answers the second question of Sec. 3.3: "how do we know which transformation to use?"

But there's a simpler way of finding the transformation $v(x)$ that doesn't involve integrals and differential equations. Let $x(t)$ be whatever function it wants – we know one transformation of $x(t)$ that's linear in time, namely the inverse function of $x(t)$, which we denote

$$x^{(-1)}(x) = t \quad (3.20)$$

$x^{(-1)}(x)$ is the transformation that pulls t out of x : I give you the value of x , you take $x^{(-1)}(x)$, and you know what t is.

So far, so good – now we know how to get t , which is of course linear in t . But that no longer tells us how fast something grows: we can't use $x^{(-1)}(x)$ as $v(x)$ because $\frac{x^{(-1)}[x(t+\Delta t)] - x^{(-1)}[x(t)]}{\Delta t} = 1$, always. So something is missing.

If we use $x^{(-1)}$ as the transformation in the general growth rate, we're in effect measuring the speed of the process on the scale of the process, which is why the answer is trivial: we will always find a growth rate of 1. The growth rate is a *conversion factor* between time measured on the standard clock (one that ticks once a second, say), and time measured on the process clock (one that advances γ units on the $v(x)$ -scale in each second).

So $x^{(-1)}$ has the right form but not the right scale. Instead, let's try the following: take the process $x(t)$ at unit rate on the standard clock. We'll denote this as $x_1(t)$. If we now take its inverse as the transformation, $v(x) = x_1^{(-1)}(x)$, it will of course produce a rate 1 if $\gamma = 1$. But if γ is something else, it will extract that something else for us!

Here's the algorithm for measuring the growth rate for a general process $x(t)$.

- Write down the process at rate 1 on the standard clock, $x_1(t)$.
- Invert it, to find the transformation $v = x_1^{(-1)}(x)$.

- Finally, evaluate the rate of change of the transformation of the process at the unknown growth rate,

$$g = \frac{x_1^{(-1)}[x(t + \Delta t)] - x_1^{(-1)}[x(t)]}{\Delta t} \quad (3.21) \quad \{\text{eq:g_inv}\}$$

The key conceptual message from this section is this: any growth process defines an appropriate functional form of a growth rate. If we measure a process with the wrong form of growth rate, we obtain something that's not stable in time. Measurements will be irreproducible or inconsistent, depending on arbitrary circumstances: the time of measurement or how long we measured for.

The key formal result is this: (Eq. 3.3) tells us there is a specific transformation of wealth that is required to state meaningfully how fast wealth grows. Equation (3.21) tells us what that transformation is.

The transformation is a linearisation. At the moment we could call it the stationarity transformation because it appropriately removes time dependence. Later – when we generalise to random growth processes – we will call it the ergodicity mapping. In the economics literature, the closest thing to this transformation is called the utility function – a term we will mostly avoid because it comes with unhelpful connotations.

3.4 Decisions in a deterministic world

{section:Decisions_in_a_d

Having clarified what a growth rate is, we can now apply our decision axiom to different situations: act so as to maximize the growth rate of your wealth. In other words, we can explore the world generated by this axiom. We will build from the ground up. First, in Sec. 3.4.1, we will look at comparing two simultaneous payments of different magnitude – which one will the model human choose who obeys our axiom? This is a sanity check: does the model human choose the bigger payment? Next, in Sec. 3.4.2 we add time – what if the model human chooses between two payments of different magnitudes that will occur at different times? This is already a far more complex situation, where the decision criterion requires knowledge of the dynamic. It will shed light on what's called “discounting.”

In the next chapter, Chap. 4, we will add fluctuations, noise, uncertainty: what if the model human doesn't know the magnitude (or time) of the payments with certainty? But for now, everything will be perfectly known.

3.4.1 Different magnitudes

{section:Different_magni

I'm off to the bank to withdraw some money for you. I offer to give you either

- (1) \$10 when I get back or
- (2) \$25 when I get back. You tell me what you prefer.

Let's see what our decision axiom says you'll do. Remember there's no uncertainty, I'm not lying to you, no one will rob me on my way from the bank *etc.*

I haven't told you how long it will take me to get to the bank, so we have to keep that general. We'll call that time interval Δt . Because we know that Δt is the same under options (1) and (2) we don't actually need to know its value to compare the growth rates for the two options. Nor do we have to know the functional form of the growth rate. In this simple case, we can work with a general $v(x)$ in (Eq. 3.3), and any growth rate will give the same answer. Let's see. Under option (1) we have

$$g^{(1)} = \frac{v(x + \$10) - v(x)}{\Delta t}, \quad (3.22)$$

and under option (2) we have

$$g^{(2)} = \frac{v(x + \$25) - v(x)}{\Delta t}. \quad (3.23)$$

To find out which growth rate is larger, we subtract $g^{(1)}$ from $g^{(2)}$

$$g^{(2)} - g^{(1)} = \frac{v(x + \$25) - v(x) - (v(x + \$10) - v(x))}{\Delta t} \quad (3.24)$$

$$= \frac{v(x + \$25) - v(x + \$10)}{\Delta t}. \quad (3.25)$$

Because $v(x)$ is monotonically increasing, any proper growth rate will be greater under option (2), and our model humans will always go for option (2). That's good – because we would have chosen option (2) if we were you, and our model reproduces this intuitive result.

More generally, our model says: of two certain payments of different sizes at the same time, choose the bigger one.

3.4.2 Different magnitudes and times: discounting

{section:Different_magnit

Let's make the decision a little harder: what if we offer you the same amounts as before, but now at different times:

- (1) \$10 in a month or
- (2) \$25 in two months?

Again, we will compute the two growth rates corresponding to options (1) and (2), and then choose the bigger one – that's how we have been programmed to behave in the world that our axiom is creating. But unlike in the previous case, the functional form of the growth rate will now be important.

Discounting under additive dynamics

Let's start with the additive growth rate, (Eq. 3.1), which is nothing but using the identity function in the general growth rate, $v(x) = x$ in (Eq. 3.3). Which payment do the model humans choose according to this rate? We've got all the parameters, so this is just a matter of substitution

$$g_a^{(1)} = \frac{x + \$10 - x}{1 \text{ month}} = \$120 \text{ p.a.}, \quad (3.26)$$

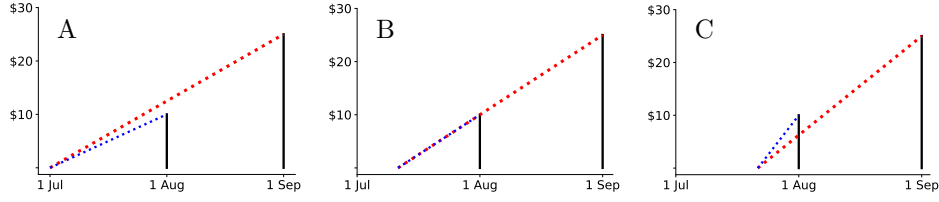


Figure 3.1: Slopes with linear vertical scales are *additive* growth rates. (A) At the beginning option (2) yields the highest additive growth rate; (B) after 1/3 of a month, the options are equally good; (C) after 2/3 of a month, preference reversal has taken place, and option (1) now yields the highest growth rate. As the first payment approaches, the associated growth rate diverges.

{fig:hyp_disc}

and

$$g_a^{(2)} = \frac{x + \$25 - x}{2 \text{ months}} = \$150 \text{ p.a..} \quad (3.27)$$

The result is clear: the decision axiom, using this growth rate, produces model humans that prefer payment (2). The additive growth rate has a unique feature: initial wealth, x cancels out. I didn't need to know your initial wealth to compute the rate! Only under additive dynamics does initial wealth not enter into the computation of the growth rate, and growth rates can be computed with knowledge of only the payouts and waiting times.

But perhaps the setup is more interesting than it seems at first glance. In the economics literature, decision-making based on additive growth rates is called “hyperbolic discounting” because this case is mathematically equivalent to discounting payments in the future with the hyperbolic function $\frac{1}{\Delta t}$.

An interesting feature of optimizing additive growth rates is what's called “preference reversal:” let's keep our example as it is, except we now let time march forward, holding fixed the moments in time when the payments are to be made. Under these conditions, there comes a time, precisely after a third of a month, when option (2) is no longer preferred, see Fig. 3.2.

You may wonder why someone might model wealth as an additive process. Here is one possibility: if wealth is mostly affected by income and expenses then it will be described by an additive dynamic. Imagine you have a monthly salary of \$1,000, and you spend \$900 every month on all your expenses. So long as any investment income, like interest payments *etc.*, is negligible, your wealth will follow (Eq. 3.4) with $\gamma = \$100$ per month.

Discounting under multiplicative dynamics

What about the exponential growth rate, with $v(x) = \ln x$ in (Eq. 3.3)? We now have growth rates

$$g_m^{(1)} = \frac{\ln(x + \$10) - \ln(x)}{1 \text{ month}}, \quad (3.28)$$

and

$$g_m^{(2)} = \frac{\ln(x + \$25) - \ln(x)}{2 \text{ months}}. \quad (3.29)$$

Curiously, which is greater depends on your initial wealth in our model world. If your wealth is \$100, then $g_m^{(1)} \approx 114\%$ p.a., and $g_m^{(2)} \approx 134\%$ p.a., wherefore you will choose option (2).

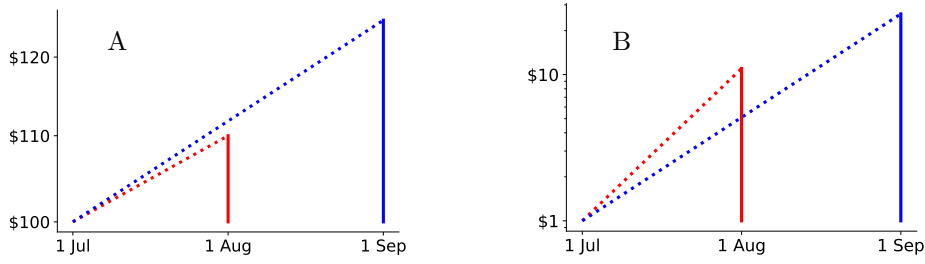


Figure 3.2: Slopes with logarithmic vertical scales. (A) If you have a lot of money (here \$100), exponential growth-rate optimization tells you to be patient and choose the later, larger, payment of \$25. (B) If you have little money (here \$1), the same criterion – exponential growth-rate optimization – tells you to get the cash as fast as possible, and choose the earlier, smaller, payment of \$10.

{fig:hyp_disc}

But if your initial wealth is \$1, then $g_m^{(1)} \approx 2,877\%$ p.a. and $g_m^{(2)} \approx 1,955\%$ p.a., and you'll choose option (1).

Notice how the poorer decision maker seems to be more impatient, despite his use of the exact same decision axiom. Using $v(x) = \ln(x)$ in (Eq. 3.3) is related to what's called “exponential discounting” in the economics literature [?]. The word “exponential” is there because we're discounting under multiplicative dynamics (which means exponential growth), and the logarithm is the inverse function of the exponential, which we need to define the appropriate growth rate, see (Eq. 3.21).

Again, let's ask why one might want to model wealth as a multiplicative process. Multiplicative processes have the property that how much you gain tomorrow is proportional to how much you currently have. Increases in wealth are proportional to wealth – this multiplicative property is virtually ubiquitous in Nature. You can imagine many reasons why it applies to wealth. For instance, wealth can be invested in interest-paying bonds. Or it can be invested in oneself: I may pay for a roof over my head, which transforms my life and earning potential from that of someone living on the streets to that of someone with a home address. Similarly, I can invest in my health and education. In Nature, multiplicativity and exponential growth occurs whenever something lives. That's because the most successful definition we have of life is “that which self-reproduces,” and self-reproduction implies multiplicative growth.

We learn: in this slightly more complex though still fully deterministic case, which option is preferable does not only depend on the options available but also on the personal circumstances of the decision maker. Both the way the decision maker thinks about wealth as a dynamical process, and his wealth influence his preferences.

Perhaps the most significant message is the richness of this problem. We're applying nothing but our simple axiom, but it forces us to choose how we think about the dynamics of our wealth, and in reality that may depend strongly on many difficult to specify circumstances. In real life payments are not just offered at some point in time, but usually in return for something – an asset or work. Depending on the specific exchange, an additive, multiplicative, or more general model will be appropriate. Such models are explored in even greater detail in [?].

Importantly, we need not resort to psychology to generate a host of behaviours, such as impatience of poorer individuals or preference reversal as time

ticks on.

Chapter 4

Decisions in a risky world

{chapter:Risky}

In the previous section we saw some interesting types of behaviour that occur in a model world generated by our decision axiom. We were able to relate these to phenomena that already have names in the more complex model worlds of classical economics, for instance “discounting” and “preference reversal.”

In the present section we will introduce randomness to our model, which will resemble situations where a decision maker is not completely sure about what the consequences of his decisions will be. We will do this in a way that’s natural from the perspective of growth rate optimization, and it will lead us to discover precisely what the meanings are in our new framework, of concepts in classical economics, such as gambles and utility theory.

We clearly have to do something about our axiom: with randomness the growth rates that are the basis of decisions in our model will also be random. We get around that problem by interpreting “growth rate” in the axiom as “time-average growth rate” when there’s randomness involved. Finding these time averages will be simple: the growth rate that’s time-invariant for a give deterministic growth process is ergodic when we introduce randomness. Its time average can therefore be computed as its expectation value, which makes this a local (in time) operation.

Decision axiom *with randomness*

People optimize the *time-average* growth rate of their wealth.

4.1 Perturbing the process

{section:Perturbing}

We begin by introducing noise into the wealth dynamic. This has to be done carefully because we want the significance of the noise to stay the same as time passes. That doesn’t necessarily mean that the amplitude of the noise – the absolute size of the typical perturbation – will stay the same. But we don’t want to have to adjust the perturbation by hand, either – we’re looking for a systematic way of perturbing the process that automatically takes into account the way in which $x(t)$ changes with time.

Without further ado, here’s the solution: introduce a constant-amplitude perturbation to something about the process that is otherwise unchanging. Of

course – you’ve guessed it – the growth rate fits the bill. In general – namely unless $x(t)$ is additive – such a perturbation will change the expectation value $\langle x(t) \rangle$, that is, in general we will have

$$\langle x \rangle(t) \neq x(t) \quad (4.1) \quad \{\text{eq:exp_changed}\}$$

and we will explore the consequences of this inequality. Incidentally, when we say $x(t)$ “is additive,” we mean that time is an additive operation. Adding some δt to time t is then equivalent to adding some δx to x .

We will follow the familiar structure, and first try out this recipe for additive dynamics, $v(x) = x$, then for multiplicative dynamics, $v(x) = \ln x$, and finally for general dynamics.

4.1.1 Perturbed additive dynamics

{section: Perturbed additi

We expect additive dynamics to be a trivial case because the additive growth rate is just the rate of change, and the stationarity transformation is the identity, $v(x) = x$. We start from deterministic additive growth, written in differential form

$$g_a = \frac{dx}{dt} = \gamma \quad (4.2)$$

then rearrange and add the perturbation. This makes sure that the dynamic significance of the perturbation doesn’t change over time: the constant growth rate becomes the ergodic growth rate. Specifically, we choose a standard Wiener perturbation.

$$dx = g_a dt + \sigma dW(t) \quad (4.3)$$

$$= \gamma dt + \sigma dW(t). \quad (4.4) \quad \{\text{eq:BMo_dx}\}$$

We integrate this (setting $x(0) = 0$ for simplicity)

$$x(t) = \int_0^t g_a ds + \sigma dW(s) \quad (4.5)$$

$$= \gamma t + \sigma W(t). \quad (4.6) \quad \{\text{eq:BM_x}\}$$

Because of additivity, in this special case we expect $x(t) = \langle x(t) \rangle$ – meaning (Eq. 4.1) is not true here. So let’s check by taking expectations

$$\langle x(t) \rangle = \langle \gamma t + \sigma W(t) \rangle \quad (4.7)$$

$$= \gamma t \quad (4.8)$$

Comparing to x , we confirm that the perturbation in this case does not change the expectation value of the process.

To recap: first, the noise in this dynamic has constant dynamic significance because it is a constant-amplitude perturbation applied to something that is unchanging in time in the deterministic case. Second, the expectation value of this dynamic is identical to the unperturbed, deterministic, case ($\sigma = 0$).

4.1.2 Perturbed multiplicative dynamics

{section: Perturbed multip

Multiplicative dynamics will be less trivial but shouldn't be too hard. The multiplicative growth rate is just rate of change of the logarithm, meaning the stationarity transformation is the logarithm, $v(x) = \ln x$. We start from deterministic multiplicative growth, written in differential form

$$g_e = \frac{d \ln x}{dt} = \gamma \quad (4.9) \quad \{\text{eq: gexp}\}$$

then, as before, rearrange and add the perturbation to the ergodic growth rate. Again this ensures that the dynamic significance of the perturbation doesn't change over time. Like in the additive case, we choose a standard Wiener perturbation.

$$d \ln x = \gamma dt + \sigma dW(t). \quad (4.10) \quad \{\text{eq: GBM_dlnx}\}$$

Just as we did to arrive at (Eq. 4.6), we have to integrate a Brownian motion. Because we're working in the ergodically transformed variable, this is a recurring theme: whatever the process $x(t)$, once it's been put through the appropriate transformation, we will end up with Brownian motion. Integrating (Eq. ??) (setting $\ln x(0) = 0$ for simplicity),

$$\ln x(t) = \int_0^t \gamma ds + \sigma dW(s) \quad (4.11)$$

$$= \gamma t + \sigma W(t). \quad (4.12) \quad \{\text{eq: GBM_lnx}\}$$

Because the stationarity transformation is no longer the identity, we now have to invert the logarithm (apply $v^{(-1)}(\cdot) = \exp(\cdot)$) to find the actual process

$$x(t) = \exp[\gamma t + \sigma W(t)]. \quad (4.13)$$

Multiplicative dynamics are not additive, meaning there is a mismatch between the additive expectation value and the multiplicative effects of time. The expectation value of an exponentiated Wiener noise is boosted by the fluctuations: the expectation value is linear, but the exponential generates disproportionately large contributions for positive values of $W(t)$. We leave it as an exercise to compute the expectation value, $\langle x(t) \rangle$, (hint: $x(t)$ is log-normally distributed) and only state the well-known result here

$$\langle x(t) \rangle = \langle \exp[\gamma t + \sigma W(t)] \rangle \quad (4.14)$$

$$= \exp \left[\left(\gamma + \frac{\sigma^2}{2} \right) t \right] \quad (4.15)$$

Comparing to x , we see that in this case, as in most cases, (Eq. 4.1) applies. If we want a multiplicative process whose expectation value is identical to its zero-noise limit, we have to include a correction. Indeed, when we consider multiplicative dynamics, we will work with this parameterization

$$d \ln x = \left(\mu - \frac{\sigma^2}{2} \right) dt + \sigma dW(t). \quad (4.16) \quad \{\text{eq: log_inc_noise}\}$$

Again: the noise in this dynamic has constant dynamic significance because it is a constant-amplitude perturbation applied to something that is unchanging

in time in the deterministic case. Thanks to the correction term $-\frac{\sigma^2}{2}$, the expectation value of this dynamic, (Eq. 4.16), is identical to the unperturbed, deterministic, case ($\sigma = 0$).

Equation (4.16) is solved for x by integrating and then exponentiating,

$$x(t) = \exp \left[\left(\mu - \frac{\sigma^2}{2} \right) t + \sigma W(t) \right]. \quad (4.17) \quad \{\text{eq:mult_x}\}$$

4.1.3 Perturbed general dynamics

{section: Perturbed_general}

The simplest, and probably most important, models to describe real-life wealth are additive, $v(x) = x$, and multiplicative, $v(x) = \ln x$. But it helps to think of those two cases as examples of something more general. In this section, we will re-do what we did for additive and multiplicative dynamics but keep things more general – within reason $x(t)$ may now follow any dynamic you like. To confirm we're really just generalising, you can go through the following steps and verify that Sec. 4.1.1 and Sec. 4.1.2 are special cases.

We find it intuitive to start with a deterministic process and add an appropriate perturbation. However, because this approach doesn't always produce the smoothest mathematics, we will return to the problem in Sec. 4.5 from a slightly different angle. That will allow us to use Itô calculus and lead to a deeper analysis, but first things first.

The starting point is now a general deterministic growth rate

$$g = \frac{dv(x)}{dt} \quad (4.18)$$

$$= \gamma. \quad (4.19) \quad \{\text{eq:growth_gen}\}$$

We re-arrange this and add the perturbation

$$dv = \gamma dt + \sigma dW. \quad (4.20) \quad \{\text{eq:dv_gen}\}$$

As before, the perturbation is guaranteed to have constant dynamic significance because it is applied to the otherwise unchanging growth rate – that's how v is defined: it's the transformation of x that grows linearly in time (so that $g = \frac{dv}{dt}$ is constant, taking the value γ in the absence of perturbations). We integrate (assuming $v(0) = 0$ for simplicity),

$$v[x(t)] = \int_0^t \gamma ds + \sigma dW \quad (4.21)$$

$$= \gamma t + \sigma W(t) \quad (4.22)$$

$$(4.23) \quad \{\text{eq:v_int}\}$$

...and apply the inverse function $v^{(-1)}$...

$$x(t) = v^{(-1)}[\gamma t + \sigma W(t)]. \quad (4.24) \quad \{\text{eq:inversion}\}$$

Equation (4.24) can be read as x being a non-linear function, $v^{(-1)}$, of a symmetrically perturbed time, $\gamma t + \sigma W(t)$. If $v^{(-1)}$ is convex, then Jensen's inequality tells us that the expectation value of this perturbed function is greater than the unperturbed case, meaning at any time

$$\langle x(t) \rangle > x(t). \quad (4.25) \quad \{\text{eq:exp_x_gen}\}$$

It can be shown that $v^{(-1)}$ is convex whenever v itself is concave (to prove this v has to be monotonically increasing, which is true by construction in our case).

The expectation value is misleading

This means, whenever the ergodicity transformation is concave, the expectation value of x is misleadingly large: it is larger than $\bar{x}$ (the unperturbed case), and grows faster (at a higher rate) than any individual trajectory of x will grow in the long run. Put simply: as regards an individual trajectory, the expectation value is pure fiction in this case. Its performance over time is not indicative of the performance of real trajectories.

We will see below that a concave ergodicity transformation (called a utility function in economics) is associated with what economists call “risk aversion.” I am deemed risk averse if I prefer the risk-free process $\bar{x}(t)$ to a risky process $x(t)$ whose expectation value is $\langle x(t) \rangle = \bar{x}(t)$. Because the expectation value is misleadingly large, it makes perfect sense to be risk averse: in the long run, the risk-free process – by Jensen’s inequality above – is guaranteed to outperform the risky one.

4.2 The appropriate growth rate is ergodic

Equation (3.3) defines deterministic dynamics by specifying a growth rate. Knowledge of the growth rate is thus knowledge of the dynamic. In the deterministic case, we defined growth rates by insisting that they not change over time. In the risky, or noisy, case, of course growth rates do change over time but – crucially – the appropriately defined growth rate changes only because of the noise and nothing else – it fluctuates from one time interval to another, but it does not systematically increase or decrease. Specifically, the way we constructed noisy dynamics, starting from the growth rate, guarantees that the noisy growth rates are ergodic: their expectation values are also their long-time averages. We will now prove this result. We will only prove it for the general case because that’s easily done, and it implies the veracity of the statement for any special case.

Proof. To avoid problems with differentiating the non-differentiable Wiener process, we begin with the expectation value of gdt ,

$$\langle gdt \rangle = \langle dv \rangle \quad (4.26)$$

$$= \langle \gamma dt + \sigma dW \rangle \quad (4.27)$$

$$= \gamma dt \quad (4.28)$$

Dividing by dt , we find the expectation value of g to be

$$\langle g \rangle = \gamma. \quad (4.29) \quad \{\text{eq:exp_value_growth_gen}\}$$

Next, we compute the time average of g . By the definition of time averages,

that's

$$\bar{g} = \lim_{T \rightarrow \infty} \frac{1}{T} \int_0^T g(t) dt \quad (4.30)$$

$$= \lim_{T \rightarrow \infty} \frac{1}{T} \int_0^T dv \quad (4.31)$$

$$= \lim_{T \rightarrow \infty} \frac{1}{T} \int_0^T [\gamma dt + \sigma dW] \quad (4.32)$$

$$= \gamma + \lim_{T \rightarrow \infty} \frac{1}{T} \sigma W(T) \quad (4.33)$$

$$= \gamma, \quad (4.34) \quad \{\text{eq:time_average_growth_g}\}$$

where the last line follows with probability one from the scaling properties of the Wiener process.

Comparing (Eq. 4.29) and (Eq. 4.34), we conclude that the expectation value and time average of g are identical for general dynamics, and hence that the growth rate, appropriately defined for a given process, is an ergodic observable. $\square$

The result shows that the two averages are identical, irrespective of whether (Eq. 4.1) applies or not. So, while the expectation value of x usually does not reflect what actually happens over time in a single system, the expectation value of the growth rate g is *guaranteed* to reflect what happens over time in a single system.

We cannot overstate the importance of this fact: *the appropriate growth rate for a stochastic wealth process is an ergodic observable*. Economic theory, historically, is based on expectation values. This is simply because when the foundations of economic theory were laid, expectation values had been invented as tools of analysis, whereas time averages had not. The ergodicity of growth rates allows us to use their expectation values, even in cases where we're actually interested in time averages. This, and its science-historical significance, will be worked out in detail in Sec. 4.4.

4.3 Ergodicity economics decision algorithm

We're now in a position to specify how humans in our model will make decisions under uncertainty. In Chap. 3 we worked out precisely how a growth rate must be defined for a given deterministic dynamic. In Sec. 4.1 we used our knowledge of such growth rates to generate consistently perturbed dynamics.

It is possible to write down dynamics that don't allow the definition of a growth rate. It is also possible to perturb a dynamic in an inconsistent way. But in these lecture notes, all deterministic dynamics will have proper growth rates, and all stochastic dynamics are obtained as consistently perturbed deterministic dynamics.

Given these constraints on the processes we consider, we can specify the algorithm our model humans will follow when making decisions under risk.

Ergodicity economics decision algorithm

When deciding which option, m^* , to choose

1. Specify the wealth dynamic, $dx(t)$, or equivalently its ergodicity transformation, $v(x)$, or equivalently the form of the relevant growth rate g ;
2. Determine the time-average growth rates, $\bar{g}^{(m)}$, either directly by averaging over time; or by invoking the ergodic property and averaging over the ensemble;
3. Choose the option, m^* , with the largest time-average growth rate.

4.4 Relation to earlier economic theories

{section:Relation_to_ear

In the previous section we presented the ergodicity economics model of decision-making: maximise the time-average growth rate of wealth. That's it; everything in ergodicity economics can be related back to this simple idea.

In the present section, we will detail precisely how our model is related to models that were developed earlier, specifically at a time before the ergodicity question had been asked. This section is thus about history: how did we model human decision making under risk before we knew we had to find appropriate growth rates?

The story is fascinating: despite the absence of appropriate tools and concepts, early researchers invented ingenious mathematical representations of human behaviour. With not very much work, it is possible to tweak these models, at least in special cases, to coincide with ergodicity economics.

The roadmap for this section is as follows. We will first introduce the concept of a *gamble* – in essence that's a little piece of a stochastic process, a random wealth change realised over some short time. Next, we will mention the *gamble problem*, namely the problem of choosing between different gambles when we have to. To connect to ergodicity economics, we point out how different ways of repeating a gamble can generate different dynamics, very similar to the types considered in Sec. 4.1. Finally, we will present the classic solution to the gamble problem, which assigns a utility $u(x)$ – non-linear in wealth – to monetary wealth. This will be summarised in the expected-utility-theory decision algorithm.

4.4.1 Gamble: random number and duration

{section:Gamble:}

One fundamental building block of mathematical decision theory is the gamble. This is a mathematical object that resembles a number of situations in real life, namely situations where we face a decision whose consequences will be purely financial and are somewhat uncertain when we make the decision. An example is buying a lottery ticket. We define the gamble mathematically as follows.

DEFINITION: Gamble

A gamble is a pair of a random variable, Q , and a duration, δt .

Q is called the payout and takes one of K (mutually exclusive) possible monetary values, $\{q_1, \dots, q_K\}$, associated with probabilities, $\{p_1, \dots, p_K\}$, where $\sum_{j=1}^K p_j = 1$. Payouts can be positive, associated with a monetary gain, or negative, associated with a loss. We order them such that $q_1 < \dots < q_K$.

In economics, the duration of a gamble is rarely discussed but it's clearly crucial information, as a trivial example shows. Let's say you get to choose between two gambles. One pays \$1 with probability 1 and takes one second to complete; the other also pays \$1 also with probability 1 but takes one year to complete. Clearly the former is more attractive. Knowing the duration will also be necessary to relate the gamble to ergodicity economics – the latter being based on time, this should come as no surprise.

The following situations may be modelled as gambles:

Example: Betting on a fair coin

Imagine betting \$10 on the toss of a fair coin. We would model this with the following payouts and probabilities:

$$q_1 = -\$10, \quad p_1 = 1/2; \quad (4.35)$$

$$q_2 = +\$10, \quad p_2 = 1/2. \quad (4.36)$$

The duration may be the time until you receive the payout, or if you participate in one coin toss every week, say, we may want to make it $\delta t = 1$ week.

Example: Playing the lottery

We can also imagine a gamble akin to a lottery, where we pay an amount, F , for a ticket which will win the jackpot, J , with probability, p . The corresponding payouts and probabilities are:

$$q_1 = -F, \quad p_1 = 1 - p; \quad (4.37)$$

$$q_2 = J - F, \quad p_2 = p. \quad (4.38)$$

Note that we deduct the ticket price, F , in the payout q_2 . The duration may be $\delta t = 1$ week.

Example: The null gamble

It is useful to introduce the null gamble, in which a payout of zero is received with certainty: $q_1 = \$0, p_1 = 1$. This represents the 'no bet' or 'do nothing' option.

As in the examples above, the duration, δt , has to be chosen appropriately. The meaning of the duration will become clearer later on – often it

is the time between two successive rounds of a gamble.

The gamble we have presented is discrete, in that the payout, Q , is a random variable with a countable (and, we usually assume, small) number of possible values. The extension to continuous random variables is natural and used frequently to model real-world scenarios where the number of possible outcomes, *e.g.* the change in a stock price over one day, is large.

This presents a natural connection to ergodicity economics: given a stochastic wealth process dx , the wealth change generated by that process over a certain time interval $[t, t + \delta t]$ is a gamble. The random variable is

$$Q = \int_t^{t+\delta t} dx, \quad (4.39) \quad \{\text{eq:gamble_from_process}\}$$

and the duration is δt .

Suppose now that you have to choose between two options that you've modelled as two gambles (possibly including the null gamble). Which should you choose, and why? This is the gamble problem, the central question of decision theory, and the basis for much of mainstream economics.

DEFINITION: The gamble problem

The gamble problem is the problem to choose between two gambles.

We stress here that the gamble alone is not enough information to answer this question. The value of a gamble – clearly – depends on more facts than probabilities, payouts, and duration. Crucially, it depends on

1. how the gamble affects our future. For instance: if we go bankrupt as a result, can we recover from that?
2. our circumstances. When you're very rich you can afford to take risks that you can't afford when you're poor.
3. our personality. Some like the thrill of gambling, others find it unpleasant.

Mainstream economics focuses on point 3. Ergodicity economics, on the other hand, focuses on points 1 and 2, where we would compute the average growth rates of the processes corresponding to the gambles. Even in the simple case of multiplicative dynamics, this requires knowledge of the individual's wealth before the gamble, $x(t)$. It also requires knowledge of the process itself. Without that we wouldn't know the form of the ergodic growth rate; we wouldn't know what to compute.

Although the gamble problem is underspecified by gambles alone, the gamble is a useful conceptual unit because it specifies that part of the model of evolving wealth that is independent of individual circumstances. It thus splits an easily observable part of the problem from information that's much harder to obtain. We can all buy the same lottery ticket with publicly specified prizes, probabilities, duration – for some of us that will be attractive, for others it won't, for a variety of reasons.

4.4.2 Repeated gamble: towards a wealth process

{section:Repeated_gamble}

We already know one touch point between ergodicity economics and the classic gamble setup: we know how to construct a gamble as the random variable corresponding to a time interval of a stochastic process – that’s (Eq. 4.39). Ergodicity economics specifies how to evaluate a wealth process, and if it’s justifiable to identify the gamble with the process, then we know how to evaluate the gamble under ergodicity economics.

We will now strengthen this connection by showing what it might mean to identify the gamble and the process. To do this we will take a gamble and construct from it a wealth process. That means we have to extend the gamble over an arbitrarily long time. A principled way of doing that (one that keeps extra assumptions to a minimum and clearly visible) is to imagine the gamble is being repeated over and over again.

Crucially, *the mode of repetition is not specified in the gamble itself*. It is the only additional assumption we have to make to arrive at a wealth process and unlock the power of ergodicity economics. We shall focus on two modes: *additive* and *multiplicative* repetition, which correspond to additive and multiplicative dynamics. Thus *the same gamble can correspond to different dynamics*.

Additive repetition

DEFINITION: Additive repetition

If a gamble is repeated additively, then a newly generated realization of the random payout, q , is added to $x(t)$ in each round. We define the change in wealth occurring over a single round as

$$\delta x(t) \equiv x(t + \delta t) - x(t). \quad (4.40) \quad \{\text{eq:DW_def}\}$$

In the additive case, we have

$$\delta x(t) = q. \quad (4.41) \quad \{\text{eq:DW_add}\}$$

In other words, under additive repetition, δx is a stationary random variable, meaning the ergodicity transformation is the identity, $v(x) = x$, and we’re in the case of additive dynamics. Starting at time, t_0 , wealth after T rounds is

$$x(t_0 + T\delta t) = x(t_0) + \sum_{\tau=1}^T q(\tau), \quad (4.42) \quad \{\text{eq:Wt_add}\}$$

where $q(\tau)$ is the realisation of the random variable in round τ . This is an evolution equation for wealth following a noisy additive dynamic. Note that $x(t_0 + T\delta t)$ is itself a random variable.

Example: Additive repetition of a \$10 bet

We return to our first example of a gamble: a \$10 bet on a coin toss. Under additive repetition, successive bets will always be \$10, regardless of how rich or poor you become. Suppose your starting wealth is $x(t_0) = \$100$.

Then, following (Eq. 4.42), your wealth after T rounds will be

$$x(t_0 + T\delta t) = \$100 + \$10k - \$10(T - k) \quad (4.43)$$

$$= \$[100 + 10(2k - T)], \quad (4.44)$$

where $0 \leq k \leq T$ is the number of tosses you've won. Note that we have assumed your wealth is allowed to go negative. If not, then the process would stop when $x < \$10$, since you would be unable to place the next \$10 bet.

Multiplicative repetition

An alternative is multiplicative repetition. In the example above, let us imagine that the first \$10 bet were viewed not as a bet of fixed monetary size, but as a fixed fraction of the starting wealth (\$100). Under multiplicative repetition, each successive bet is for the same fraction of wealth which, in general, will be a different monetary amount.

The formalism is as follows.

DEFINITION: Multiplicative repetition

The payout, q_j , in the first round is expressed as a random wealth multiplier,

$$r_j \equiv \frac{x(t_0) + q_j}{x(t_0)}. \quad (4.45) \quad \{\text{eq:R_def}\}$$

This multiplier is another random variable, and multiplicative repetition means drawing a new instance of it every δt and multiplying wealth $x(t)$ accordingly. Wealth after T rounds of the multiplicatively repeated gamble is

$$x(t_0 + T\delta t) = x(t_0) \prod_{\tau=1}^T r(\tau), \quad (4.46)$$

where $r(\tau)$ is the realisation of the random multiplier in round τ . The ergodicity transformation is now $v(x) = \ln x$, by which we mean that logarithmic wealth changes, $\delta \ln x$, are stationary. The exponential growth rate is ergodic, and its expectation value

$$\frac{\langle \delta \ln x \rangle}{\delta t} \quad (4.47)$$

is the time-average growth rate under this mode of repetition.

Example: Multiplicative repetition

The \$10 bet on a coin toss is now re-expressed as a bet of a fixed fraction of wealth at the start of each round. Following (Eq. 4.45), the random

multiplier, r , has two possible outcomes:

$$r_1 = \frac{\$100 - \$10}{\$100} = 0.9, \quad p_1 = 1/2; \quad (4.48)$$

$$r_2 = \frac{\$100 + \$10}{\$100} = 1.1, \quad p_2 = 1/2. \quad (4.49)$$

The wealth after T rounds is, therefore,

$$x(t_0 + T\delta t) = \$100 (1.1)^k (0.9)^{T-k}, \quad (4.50)$$

where $0 \leq k \leq T$ is the number of winning tosses. In this example there is no need to invoke a ‘no bankruptcy’ condition, since we can lose no more than 10% of our wealth in each round.

We have defined a gamble and clarified how it is related to ergodicity economics. It can be derived from a wealth dynamic; conversely a wealth dynamic can be constructed from a gamble provided we know how to repeat the gamble. But early decision theory didn’t use the concepts of repetition, time averages, or ergodicity transformations. In the next section we will present how the gamble problem was addressed before the advent of ergodicity economics, and precisely how this earlier treatment is related to ours.

4.4.3 Expected wealth and expected utility

In ergodicity economics, we solve the gamble problem by maximizing the time (or ensemble) average of the ergodic growth rate for the process we consider the gamble to be part of.

But until recently, this was not how economists treated the gamble problem. It’s instructive to mention two earlier criteria for gamble evaluation: the expected wealth change, invented around 1654; and the expected utility change, invented around 1738. Despite the conceptual error embedded in these criteria – confusing expectation values with time averages – they are easy to express in terms of time averages and ergodicity economics. In later developments, such as extensions of expected-utility theory beginning in the 1930s, and of prospect theory in the 1970s, the original error is compounded, and it is unclear how to ascribe physical meaning to the resulting models.

Expected wealth

The first gamble evaluation criterion, which emerged in the early days of probability theory in the 17th century, was this:

Expected wealth decision algorithm

1. Specify $Q^{(m)}$ for the gambles offered;
2. Determine the change of expected wealth induced by each gamble,

$$\langle \delta x \rangle^{(m)} = \langle Q^{(m)} \rangle; \quad (4.51) \quad \{\text{eq:EW_criterion}\}$$

3. Choose the gamble, m^* , with the largest $\langle \delta x \rangle^{(m)}$.

The history of this criterion is linked to finding what was perceived as a fair value of an uncertain prospect. Say we're playing a game of chance, for some amount of money in a pot, maybe we're rolling dice, best out of three. If you're currently ahead, say after two throws, and someone wants to take over your position, it may be fair to sell your position for the expectation value of your winnings. One reason this criterion makes some sense is conservation: the sum of the expected winnings of all players is precisely the pot, so buying everyone out of his position is equivalent to buying the pot, as perhaps it should be.

Connecting this back to ergodicity economics: under what conditions would we evaluate a gamble by computing $\langle \delta x \rangle$? The answer is: under additive dynamics. The ergodicity economics maximand is then $\langle g_a \rangle = \frac{\langle \delta x \rangle}{\delta t}$. Apart from the δt in the denominator (which cancels if we assume it's the same for all considered gambles), this is expected-wealth maximisation. We conclude that expected-wealth maximisation is equivalent to ergodicity economics under the assumption of additive wealth dynamics: one of the simplest wealth models we can write down.

There is even a good reason why additive dynamics may have been assumed in the early developments. Let's Taylor-expand the time average of the general ergodic growth rate (which we may write as an ensemble average because of ergodicity)

$$\langle g \rangle = \frac{\langle \delta v(x) \rangle}{\delta t} \quad (4.52)$$

$$= \frac{1}{\delta t} \left\langle \left[\frac{dv}{dx} \delta x + \frac{1}{2} \frac{d^2 v}{dx^2} \delta x^2 + \dots \right] \right\rangle \quad (4.53)$$

$$= \frac{1}{\delta t} \left[\frac{dv}{dx} \langle \delta x \rangle + \frac{1}{2} \frac{d^2 v}{dx^2} \langle \delta x^2 \rangle + \dots \right] \quad (4.54)$$

The last line follows from the fact that the derivatives of v are known functions, evaluated at the known wealth x before the gamble takes place. In other words, they are just constants and can be taken out of the expectation operator, $\langle \cdot \rangle$. If δx is small, keeping only the first term in the expansion is often a valid approximation, in which case we will call the dynamic “linearisable,” and the expression becomes

$$\approx \frac{\langle \delta x \rangle}{\delta t} \frac{dv}{dx}, \quad (4.55)$$

which is proportional to the expected wealth change. In behavioural terms, being proportional means yielding the same ranking of any set of gambles as just the expected wealth $\langle \delta x \rangle$.

Summary: the first formal decision theory is expected wealth maximization. This theory is equivalent to ergodicity economics under additive dynamics, and additive dynamics are equivalent to any linearisable dynamic in the small-stakes limit.

Expected utility

In the language of economics, the expected-wealth paradigm treats humans as ‘risk neutral’, *i.e.* they have no preference between gambles whose expected changes in wealth are identical (over the same time interval). For example, people would be indifferent to either keeping what they have (the null gamble), or tossing a coin to win or lose \$1,000.

At least since 1713 [37, p. 402], this has been known to be a flawed model. Nicolas Bernoulli pointed out that gambles can be constructed – at least in theory – whose expected wealth change does not exist. What then, would this criterion mean? Moreover, expected wealth maximisation does not always accord well with observed behaviour. For instance, anyone who buys an insurance contract prefers the certain loss of the insurance fee to a random loss of smaller expectation value. Such a person does not maximise expected wealth, but insurance contracts have been signed since Phoenecian times.

In 1738 Daniel Bernoulli put forward a new model of human behaviour that addressed some of the empirical failures of expected wealth maximisation¹.

Bernoulli noted that the value to an individual of a possible change in wealth depends on how much wealth the individual already has and on his psychological attitude to taking risks. In other words, people do not treat equal amounts of extra money equally. This makes intuitive sense: an extra \$10 is much less significant to a rich man than to a pauper for whom it represents a full belly; an inveterate gambler has a different attitude to risking \$100 on the spin of a roulette wheel than a prudent saver, their wealths being equal.

In 1738 Bernoulli [6], after correspondence with Cramer, devised the ‘expected-utility paradigm’ to model these considerations. He observed that money may not translate linearly into usefulness and assigned to an individual an idiosyncratic utility function, $u(x)$, that maps his wealth, x , into usefulness, u . He claimed that this was the true quantity whose expected change, $\langle \delta u(x) \rangle$, is maximised in a choice between gambles.

This is the axiom of utility theory. It leads to yet another decision algorithm.

¹Bernoulli’s original paper contains an error. The theory as we present it here is a corrected version that can be found in [25], see also [41] for a discussion of the error

Expected-utility decision algorithm

1. Specify $Q^{(m)}$ for the gambles offered;
2. Specify the individual's idiosyncratic utility function, $u(x)$, which maps his wealth to his utility;
3. Determine the change of expected utility induced by the gamble,

$$\langle \delta u \rangle = \left\langle u \left(x + q^{(m)} \right) \right\rangle - u(x); \quad (4.56) \quad \{\text{eq:EUT_criterion}\}$$

4. Choose the gamble, m^* , with the largest $\langle \delta u \rangle^{(m)}$.

Let's ask the same question as for the expected-wealth paradigm: under what conditions would ergodicity economics maximise the quantity in (Eq. 4.56)? The utility function $u(x)$ is defined – somewhat circularly, as *e.g.* von Neumann and Morgenstern pointed out [63, p. 28] – as the object whose expected changes are maximised by a person. Under ergodicity economics, what's being maximised is the expected (rate of) change of the ergodicity transformation, $v(x)$.

Mapping: expected-utility theory $\Leftrightarrow$ ergodicity economics

We conclude that expected-utility theory is equivalent to ergodicity economics if gambles of equal duration, δt are considered and *if the utility function, u , coincides with the ergodicity transformation v .*

But it gets even better. Bernoulli did not just write down a general function $u(x)$ but argued that the logarithm is a plausible candidate for this function. People, he claimed, tend to behave as if they were optimising expected changes in the logarithm of wealth. Quite why that was the case, he didn't know. We do: using the logarithm as the ergodicity transformation, *i.e.* equating $u = v$, we find that Bernoulli's observation can be rephrased: people commonly maximise the time average of the ergodic growth rate under multiplicative dynamics. That's the second important wealth model we identified!

This is quite an astonishing correspondence, and we believe it is not coincidental. In the 18th century, researchers discovered elements of the mathematical structure of ergodicity economics. Just by careful observation – the appropriate mathematical tools had yet to be invented. Because mathematical concepts were immature at the time, a nomenclature emerged (“utility,” “risk preferences” *etc.*) that seems quaint from today's conceptual context.

Summary: historically, the second formal decision theory is expected utility maximisation. This theory is equivalent to ergodicity economics if the utility function is the ergodicity transformation. Each dynamic thus has a corresponding growth-optimal utility function. The special case treated by Bernoulli in detail – logarithmic utility – is equivalent to ergodicity economics under multiplicative dynamics. Expected wealth maximisation is a special case of expected-utility maximisation, namely using a linear utility function (corresponding to ergodicity economics under additive dynamics).

4.5 From ergodicity transformation to dynamic and back – Itô

{section:From_growth}

In this section we will use Itô calculus to consolidate and extend our results a little. In addition to what we already know (converting ergodicity transformations into wealth processes), we will arrive at a recipe for going the other way: converting wealth processes, dx , into ergodicity transformations v . This is an important part of the connection to expected-utility theory: give us the growth process, and we will tell you which utility function will outperform all others in the long run.

Furthermore, we will arrive at a set of conditions for these mapping to be possible. We illustrate the procedure with one example for each direction: we derive the growth process that corresponds to a square-root ergodicity transformation; and we derive the utility function (which will be exponential) that corresponds to a curious-looking dynamic. In other words, we go explicitly beyond linear and logarithmic ergodicity transformations.

4.5.1 Itô setup

{section:Ito_setup}

We have seen that specifying a process $x(t)$ is equivalent to specifying an ergodicity transformation $v(x)$. When we mapped one onto the other, we had to invert v . This tells us one restriction: v has to be invertible. In this section, we will do the mapping again, but using Itô calculus, which will allow us to say more about the restrictions on v and x .

We will also revisit the inequality (Eq. 4.25) and discover ageing – a dependence on t in its magnitude.

We start with the self-imposed restriction that wealth dynamics are an Itô process, which you may remember from (Eq. 2.72) in Sec. 2.12. We will further restrict ourselves to coefficient functions $a_x(x)$ and $b_x(x)$ without explicit t dependence, meaning wealth will follow

$$dx = a_x(x)dt + b_x(x)dW. \quad (4.57) \quad \{\text{eq:dx_Ito}\}$$

In this phrasing, we can implement additive dynamics by setting $a_x = \mu$ and $b_x = \sigma$ as constants (Brownian motion, (Eq. 4.4)). We can also choose multiplicative dynamics, with $a_x = \mu x$ and $b_x = \sigma x$ (geometric Brownian motion, (Eq. 4.10)). So Itô processes include the most important models we've already seen, and many others.

Our aim is to find pairs of processes dx and ergodicity transformations $v(x)$, meaning we're after a connection between an Itô process and a function of an Itô process. That's exactly what Itô's formula is for. By construction, v follows a Brownian motion with drift, which is another Itô process, namely

$$dv = a_v(v)dt + b_v(v)dW \quad (4.58)$$

$$= \gamma dt + \sigma dW \quad (4.59) \quad \{\text{eq:dv}\}$$

4.5.2 From ergodicity transformation to wealth process

Previously, we wrote x in terms of the inverse function of v as $x = v^{(-1)}(\gamma t + \sigma W(t))$. Itô's formula allows us to write dx in terms of the coefficient functions

of dv in (Eq. 4.59) as follows

$$dx = \underbrace{\left(\frac{\partial x}{\partial t} + a_v \frac{\partial x}{\partial v} + \frac{1}{2} b_v^2 \frac{\partial^2 x}{\partial v^2} \right)}_{a_x(x)} dt + \underbrace{b_v \frac{\partial x}{\partial v}}_{b_x(x)} dW \quad (4.60) \quad \{\text{eq: dx}\}$$

This expression involves partial derivatives of $x(v)$, *i.e.* of the inverse function of v (of course, $x(v) = v^{-1}(v)$). This confirms the constraint we already knew: v has to be invertible. Another constraint – clearly – is that $x(v)$ has to be twice-differentiable.

We have thus shown that

Invertible ergodicity transformations (utility functions) have dynamic interpretations

For any invertible ergodicity transformation $v(x)$ a class of corresponding wealth processes dx can be obtained such that the rate of change (*i.e.* the additive growth rate) in the expectation value of net changes in utility is the time-average growth rate of wealth.

Optimising the expected changes of the ergodicity transformation, $\langle \Delta v \rangle$, is equivalent to optimising time-average wealth growth for the corresponding wealth process, $\bar{g}(x)$.

We caution that it will be impossible for some processes $x(t)$ to find a $v(x)$ that satisfies (Eq. 4.59). In this case we cannot interpret expected utility theory dynamically, and such processes are likely to be pathological.

In the language of utility theory, every invertible utility function is an encoding of a wealth dynamic. Under that dynamic, behaving according to the corresponding utility function is optimal over time. The dynamic arises as utility – or rather: v – performs Brownian motion, and wealth is the transformation $v^{(-1)}(v(t))$.

Equation (4.60), creates the now familiar pairs of ergodicity transformations (utility functions) $v(x)$ and dynamics dx . Below we state the two familiar examples and work out a third one to illustrate the generality and ease of using Itô calculus.

Examples

- The linear ergodicity transformation (utility function) corresponds to additive wealth dynamics (Brownian motion),

$$v(x) = x \quad \leftrightarrow \quad dx = a_v dt + b_v dW, \quad (4.61)$$

as is easily verified by substituting $x(v) = v$ in (Eq. 4.60).

- The logarithmic ergodicity transformation (utility function) corresponds to multiplicative wealth dynamics (geometric Brownian motion),

$$v(x) = \ln(x) \quad \leftrightarrow \quad dx = x \left[\left(a_v + \frac{1}{2} b_v^2 \right) dt + b_v dW \right]. \quad (4.62)$$

- To demonstrate the generality of our procedure, we carry it out for another special case that is historically important. The first utility function ever to be suggested was the square-root function $u(x) = x^{1/2}$, by Cramer in a 1728 letter to Daniel Bernoulli, partially reproduced in [6]. What would be the dynamic under which it is optimal to behave according to this utility function? In other words, what dx corresponds to the ergodicity transformation $v(x) = x^{1/2}$? This case will display ageing – a new phenomenon that we haven’t encountered yet. So we’ll got through it slowly.

The square-root function is invertible, namely $x(v) = v^2$, and this inverse is twice-differentiable. So we can use (Eq. 4.60). In the next three lines we

- substitute v^2 for $x(v)$ in (Eq. 4.60)
- carry out the differentiations
- substitute $x^{1/2}$ for v .

$$dx = \left(a_v \frac{\partial v^2}{\partial v} + \frac{1}{2} b_v^2 \frac{\partial^2 v^2}{\partial v^2} \right) dt + b_v \frac{\partial v^2}{\partial v} dW \quad (4.63)$$

$$= (2\gamma v + \sigma^2) dt + 2\sigma v dW \quad (4.64) \quad \{\text{eq:dx_in_v}\}$$

$$= (2\gamma x^{1/2} + \sigma^2) dt + 2\sigma x^{1/2} dW. \quad (4.65)$$

We’ve thus established the third mapping in our collection:

$$v(x) = x^{1/2} \leftrightarrow dx = \left(2a_v x^{1/2} + b_v^2 \right) dt + 2b_v x^{1/2} dW. \quad (4.66) \quad \{\text{eq:dx_2}\}$$

The new phenomenon we’ve mentioned – ageing – is observed when we compare the time average growth rate and that of the expectation value. By construction, the time-average growth rate is $\bar{g}(x) = \frac{\langle \Delta v \rangle}{\Delta t} = \gamma$.

To compare this to the growth rate of the expectation value, we compute $\langle x \rangle$ as follows. Take the expectation value of (Eq. 4.64)

$$\langle dx \rangle = (2\gamma \langle v \rangle + \sigma^2) dt \quad (4.67)$$

$$= (2\gamma t + \sigma^2) dt \quad (4.68)$$

This can be integrated to yield

$$\langle x \rangle = \gamma t^2 + \sigma^2 t + \langle x(0) \rangle \quad (4.69)$$

and we find its growth rate by differentiating,

$$\frac{dv \langle x \rangle}{dt} = \frac{\partial v(\langle x \rangle)}{\partial \langle x \rangle} \times \frac{d\langle x \rangle}{dt} \quad (4.70)$$

$$= \frac{1}{2} \langle x \rangle^{-1/2} (2\gamma t + \sigma^2) \quad (4.71)$$

$$= \frac{(\gamma^2 t + \frac{1}{2} \sigma^2)}{(\gamma^2 t^2 + \sigma^2 t)^{1/2}} \quad (4.72)$$

The first thing to notice here is that the expectation-value growth rate, appropriately defined, is not constant but depends on when we measure

it! This is what is meant by “ageing:” the system, the dynamic itself, changes with time.

The second thing to notice is the asymptotic value. In the long run, the growth rate of the expectation value, for the Cramer dynamic converges to the ergodic growth rate,

$$\lim_{t \rightarrow \infty} \frac{dv(\langle x \rangle)}{dt} = \gamma. \quad (4.73)$$

It’s quite subtle. While Jensen’s inequality still holds,

$$\frac{dv(\langle x \rangle)}{dt} > \frac{d\langle v(x) \rangle}{dt}, \quad (4.74)$$

the magnitude of this inequality vanishes in the long-time limit. We leave the discussion of the Cramer dynamic here. It shows that ergodicity economics is full of open interesting avenues. Try out your own dynamics, and discover new phenomena!

History: Bounded utility functions in mainstream economics

Curiously, a celebrated but erroneous paper by Karl Menger [36] “proved” that all utility functions must be bounded (the proof is simply wrong). Boundedness makes utility functions non-invertible and precludes the developments we present here. Influential economists lauded Menger’s paper, including Paul Samuelson [52, p. 49] who called it “a modern classic that [...] stands above all criticism.” This is one reason why mainstream economics has failed to use the optimisation of wealth growth over time to understand human behavior – a criterion we consider extremely simple and natural. A discussion of Menger’s precise errors can be found in [48, p. 7]. Although mainstream economics still considers boundedness of utility to be formally required, it is such an awkward restriction that John Campbell noted recently [12] that “this requirement is routinely ignored.”

4.5.3 From wealth process to ergodicity transformation

{section:From_wealth}

We now ask under what circumstances the procedure in (Eq. 4.60) can be inverted: when can we find an ergodicity transformation for a given dynamic? Where this is possible, optimisation over time can be represented by optimisation of expected-utility changes.

We ask whether a given dynamic can be mapped into a $v(x)$ that follows Brownian motion, (Eq. 4.59).

In Sec. 4.5.1 we restricted ourselves to wealth following an Itô process, so that (Eq. 4.57) applies, with $a_x(x)$ and $b_x(x)$ as arbitrary functions of x . For this dynamic to translate into a Brownian motion for $v(x)$, (Eq. 4.59) must be satisfied. The tricky part here – the part that gives us constraints – is that the coefficients a_v and b_v are constants, namely γ and σ . Let’s write this in an equation

$$dv = \underbrace{\left(a_x(x) \frac{\partial v}{\partial x} + \frac{1}{2} b_x^2(x) \frac{\partial^2 v}{\partial x^2} \right)}_{a_v} dt + \underbrace{b_x(x) \frac{\partial v}{\partial x}}_{b_v} dW. \quad (4.75) \quad \{\text{eq:du}_2\}$$

To avoid clutter, let's use Lagrange notation, namely a dash $'$ to denote a derivative. Explicitly, we arrive at two equations for the coefficients

$$a_v = a_x(x)v' + \frac{1}{2}b_x^2(x)v'' \quad (4.76) \quad \{\text{eq:A}\}$$

and

$$b_v = b_x(x)v'. \quad (4.77) \quad \{\text{eq:b_u}\}$$

Differentiating (Eq. 4.77), it follows that

$$v''(x) = -\frac{b_v b_x'(x)}{b_x^2(x)}. \quad (4.78)$$

Substituting in (Eq. 4.76) for v' and v'' and solving for $a_x(x)$ we find the drift term as a function of the noise term,

$$a_x(x) = \frac{a_v}{b_v}b_x(x) + \frac{1}{2}b_x(x)b_x'(x). \quad (4.79) \quad \{\text{eq:consistency}\}$$

This equation is a consistency check of the wealth dynamic dx . An ergodicity transformation exists if and only if the coefficient functions $a_x(x)$ and $b_x(x)$ satisfy (Eq. 4.79). We do not need to construct the mapping explicitly to know whether a pair of drift term and noise term is consistent or not.

But once we know that it exists we can construct it by substituting for $b_x(x)$ in (Eq. 4.76). This yields a differential equation for v

$$a_v = a_x(x)v' + \frac{b_v^2}{2v'^2}v'' \quad (4.80)$$

or

$$0 = -a_v v'^2 + a_x(x)v'^3 + \frac{b_v^2}{2}v''. \quad (4.81)$$

Overall, then the triplet noise term, drift term, ergodicity transformation is interdependent. Given a noise term we can find consistent drift terms, and given a drift term we find a consistency condition (differential equation) for the ergodicity transformation. These arguments may seem a little esoteric when first encountered, using bits and pieces from different fields of mathematics. But they constitute the actual physical story behind the fascinating history of decision theory. To prevent the discussion from getting too dry, let's illustrate the procedure with an example.

Example: a curious-looking dynamic

We will work with the following wealth dynamic

$$dx = \left(\frac{a_v}{b_v}e^{-x} - \frac{1}{2}e^{-2x} \right) dt + e^{-x} dW. \quad (4.82) \quad \{\text{eq:test_dyn}\}$$

A typical trajectory of (Eq. 4.82) is shown in Fig. 4.1.

We will check whether an ergodicity mapping exists for it, find that to be the case, then construct the ergodicity transformation explicitly, and discuss some of its properties.

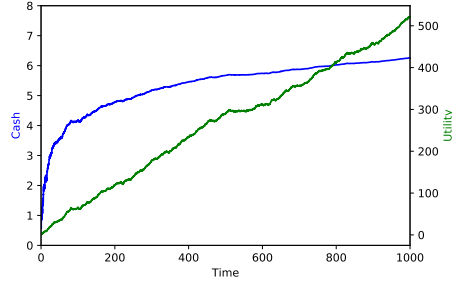


Figure 4.1: Typical trajectory $x(t)$ of the wealth dynamic (Eq. 4.82), with parameter values $a_v = 1/2$ and $b_v = 1$, and the corresponding Brownian motion $v(t)$. Note that the fluctuations in $x(t)$ become smaller for larger wealth.

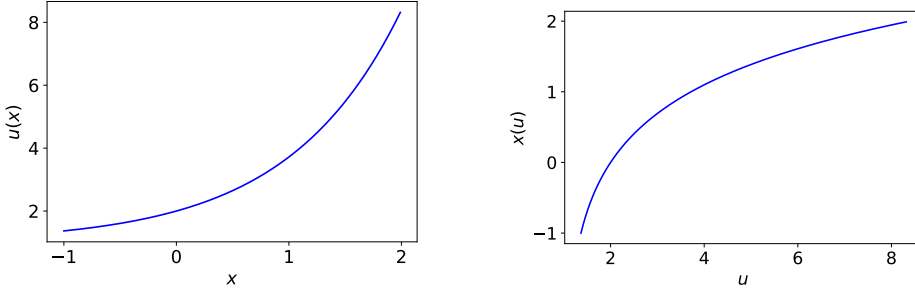


Figure 4.2: The exponential ergodicity transformation (or utility) $v(x)$, (Eq. 4.86) with $b_v = 1$ and $C = 1$, is monotonic and unbounded and therefore invertible. Left panel: $v(x)$. Right panel: inverse $x(v)$.

Check consistency: In (Eq. 4.82), we have $a_x(x) = \frac{a_v}{b_v} e^{-x} - \frac{1}{2} e^{-2x}$ and $b_x(x) = e^{-x}$. Equation (4.79) imposes conditions on the drift term $a_x(x)$ in terms of the noise term $b_x(x)$. Substituting in (Eq. 4.79) reveals that the consistency condition is satisfied by the dynamic in (Eq. 4.82).

Construct the ergodicity transformation: Because (Eq. 4.82) is internally consistent, it is possible to derive the corresponding ergodicity transformation. Equation (4.77) is a first-order ordinary differential equation for $v(x)$

$$v'(x) = \frac{b_v}{b_x(x)}, \quad (4.83) \quad \{\text{eq:diff_eq_u}\}$$

which can be integrated to

$$v(x) = \int_0^x d\tilde{x} \frac{b_v}{b_x(\tilde{x})} + C, \quad (4.84)$$

with C an arbitrary constant of integration.

Substituting for $b_x(x)$ from (Eq. 4.82), (Eq. 4.83) becomes

$$v'(x) = b_v e^x, \quad (4.85)$$

which is easily integrated to

$$v(x) = b_v e^x + C, \quad (4.86) \quad \{\text{eq:test_dyn_u}\}$$

plotted in Fig. 4.2.

Discussion: This exponential ergodicity transformation is monotonic and therefore invertible – we knew that because the consistency condition is satisfied.

The ergodicity transformation is convex. From the perspective of expected-utility theory (where the ergodicity transformation is called a utility function), an individual behaving optimally according to this function would be labelled “risk-seeking.” This behavior would commonly be observed under the dynamic (Eq. 4.82) because Jensen’s inequality in (Eq. 4.25) now points the other way from what we had in all previous examples: the expectation value of x *understates* what typically happens to a single trajectory over time.

Another new phenomenon: under the dynamic (Eq. 4.82), optimal behaviour is “risk-seeking,” in the sense that it will lead to faster wealth growth than risk-averse behaviour. This dynamic has the feature that fluctuations in wealth become smaller as wealth grows. High wealth is therefore sticky – an individual will quickly fluctuate out of low wealth and into higher wealth. It will then tend to stay there.

Once more, we see the conceptual difference to mainstream economics and utility theory: from the perspective of ergodicity economics what’s optimal is determined by the dynamic, not by the individual. Of course the individual may choose not to behave optimally.

We end here the abstract discussion of dynamics and ergodicity transformations. Much remains to be discovered in this space, but we’d like to get on to real-world applications. Before we do that, we can’t help ourselves but to hint at another instructive example, which we invite the reader to explore. Auto-elastic ergodicity transformations of the form

$$v(x) = \frac{x^{1-\eta} - 1}{1 - \eta}. \quad (4.87)$$

These can be concave or convex, depending on the parameter η . For $\eta > 1$ they are bounded, and that leads to curious problems. Because $v(x)$ is bounded, no matter how large x becomes, some values of $v(x)$ cannot occur. At the same time, our framework assumes that $v(t)$ performs a Brownian motion with drift and that it will eventually reach any value. All of this creates an internal conflict that corresponds to finite-time singularities in wealth. In other words: assuming bounded utility functions, when translated into dynamic terms, amounts to assuming finite-time singularities in wealth. Beyond these singularities, wealth becomes a complex number – mathematically, this is fun, but of course physical realism is then lost. It is interesting, of course, that Menger, with his bounded utility functions, and those who endorsed him inadvertently argued for all of economics to take place in this non-physical realm outside the real numbers. A formal disaster.

Chapter 5

Decisions in the real world

{chapter:Real}

In the previous chapter we established the conceptual foundations of ergodicity economics, including the formal relationship with expected utility theory, the dominant thinking in economics prior to ergodicity economics. We used most of the tools introduced in Chap. 2, so from here on we can forge ahead and try out some applications. By applications we mean both theoretical applications of our new conceptual toolkit, which will allow us to solve theoretical problems or elicit new aspects of the problems, and we also mean empirical work with data collected in the lab or in the wild.

In the present chapter we will discuss the early unsystematic observations surrounding the St. Petersburg paradox that led to the development of expected utility theory; a lab experiment carried out in Copenhagen that puts ergodicity economics to the test; and we will discuss the surprisingly foundational theoretical problem of insurance: why and when do people buy insurance, and how much do they pay for it?

5.1 The St. Petersburg paradox

{section:SPP}

For a bit of history and cultural background, we felt it was a good idea to discuss the St. Petersburg paradox, from the perspective of ergodicity economics.

The motivating problem was suggested by Nicolaus Bernoulli in 1713 in his correspondence with Montmort [37]. It involves a hypothetical lottery for which the rate of change of expected wealth diverges for any finite ticket price. The expected-wealth paradigm would predict, therefore, that people are prepared to pay any price to enter the lottery. However, it seems that researchers in the early 18th century conducted rudimentary informal experiments and presented the gamble to friends or strangers, asking them how much they'd be willing to wager. In other words, the expected-wealth paradigm was put to a real-world test. And it failed unequivocally: no one wanted to wager more than a few dollars. This is the first well-documented example of the inadequacy of the expected-wealth paradigm as a model of human decision-making. It was the primary motivation for Daniel Bernoulli's and Cramer's development of an alternative model: the expected-utility paradigm [6].

It's useful to separate two aspects of the paradox. First, at a superficial level, the divergent rate of change of expected wealth, irrespective of the fee, may be

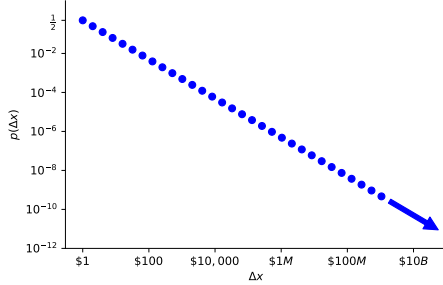


Figure 5.1: The St. Petersburg gamble, at zero fee, defines (q_j, p_j) for all $j \in \mathbb{N}^+$, where $q_j = \$2^{j-1}$, and $p_j(q_j) = \frac{\$1}{2} q_j^{-1}$ decays as a power law with exponent -1 (the straight line continues indefinitely). This implies a non-existent expectation value of Q . Since the expectation value is linear, subtracting any finite fee F from the possible prizes still leaves the expectation value of net wealth changes positively divergent.

{SC02}

{fig:SP_power_law}

taken to translate into the real-world prediction that people will be willing to pay very high prices for a ticket. That's not observed, and the prediction is *quantitatively wrong*. Second, on a deeper level, the entire expected-wealth paradigm fails because it produces no prediction for the maximum fee. This quantity diverges under this paradigm – a good place to remember that an infinity as a model output is an error message. We will consider the paradox to be the non-existence of a maximum fee. As we address this core issue, it will turn out that the quantitative values for this fee are realistic.

In some sense it is a pity that this deliberately provocative and unrealistic lottery has played such an important role in the development of classical decision theory. As we've seen, it is quite unnecessary to invent a gamble with a divergent change in expected wealth to expose the flaws in the expected-wealth paradigm. The presence of infinities in the problem and its variously proposed solutions has caused much confusion, and permits objections on the grounds of physical impossibility. These objections don't much advance decision theory: they address only the gamble and not the decision paradigm. Nevertheless, the paradox is an indelible part not only of history but also of the current debate [43], and so we recount it here. We'll start by defining the lottery.

The St. Petersburg lottery requires the player to pay a fee F and offers a random prize whose realisations are power-law distributed, with exponent -1 . Neither the duration nor the dynamic are specified in the setup.

Phrased as a gamble, we have the random payout Q with realisations

$$\forall j \in \mathbb{N}^+, q_j = \$2^{j-1} - F \quad (5.1) \quad \{\text{eq:SPQ}\}$$

and corresponding probabilities

$$p_j = \left(\frac{1}{2}\right)^j. \quad (5.2) \quad \{\text{eq:SPP}\}$$

It's easiest to see the power-law distribution, which causes the divergent first moment, by plotting q_j vs. p_j on double-logarithmic scales, shown in Fig. 5.1 where we've chosen to set the fee $F = 0$.

The problem was considered a paradox because it was posed at a time when the expected-wealth paradigm reigned supreme. This means researchers at the time computed – or failed to compute – the expected wealth change $\langle \Delta x \rangle = \langle Q \rangle$ of the lottery. The idea was to compute this as a function of the fee, and find

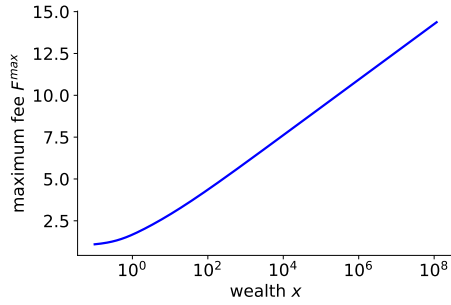


Figure 5.2: Maximum fee $F^{\max}$ which a player is willing to pay *vs.* the player's wealth x , according to ergodicity economics, assuming multiplicative dynamics. For a given wealth x , this is the value of F where the time-average of the ergodic growth rate (Eq. 5.3) is zero. This resolves the 300-year old paradox because for any finite x a finite $F^{\max}$ exists. Adapted from [43].

{SC03}

{fig:gbar_zero}

the maximum fee people would be willing to pay, that is find $F = F^{\max}$ s.t. $\langle Q \rangle = 0$.

With the random payout Q specified by (Eq. 5.1) and (Eq. 5.2), no such fee exists. The expected wealth change diverges positively for any finite fee. The expected-wealth paradigm thus fails, and this precipitated the development of expected-utility theory in the 18th century.

We will resolve the paradox here via ergodicity economics. This is done by specifying the wealth dynamic which we assume the wealth of a potential player to be subjected to. This is a guess: we would look at the circumstances of a person and conclude something like “in the relevant regime his wealth is roughly multiplicative” or whatever dynamic we guess. We will then compute the ergodic growth rate corresponding to participation in the lottery, and find the value $F^{\max}$ at which this growth rate becomes zero. Since dynamics are generally non-additive, ergodicity economics, unlike the expected-wealth paradigm, does not run aground – for many realistic dynamics such a fee exists.

Without further ado, here's the solution. We assume multiplicative wealth dynamics whose ergodic growth rate is $\frac{\Delta \ln x}{\Delta t}$. The dynamic is not specified in the problem statement, but neither is a utility function or anything else that would allow a solution: the problem as stated is underspecified.

Now compute the time average of this growth rate, which we can do by invoking ergodicity and computing its expectation value.

$$\frac{\langle \Delta \ln x \rangle}{\Delta t} = \frac{\langle \ln(x + q - F) \rangle - \ln(x)}{\Delta t}. \quad (5.3) \quad \{\text{eq:SPg}\}$$

While there's no simple closed-form solution for $F^{\max}$, it's not hard to show that this expression converges to something finite for sufficiently large fees. It also shows us that the time-average growth depends on the player's wealth. A rich player will be willing to pay a higher fee (though not an arbitrarily large fee) than a poor player.

In Fig. 5.2 we plot $F^{\max}$ (where (Eq. 5.3) is zero) as a function of the player's wealth. These values were obtained numerically. In the inset we assume a fee of \$2 and plot the time-average growth rate (Eq. 5.3) as a function of the player's wealth.

Treatments based on multiplicative repetition have appeared sporadically in the literature, starting with Whitworth in 1870 [?, App. IV].¹ It is related

¹Whitworth was dismissive of early utility theory: “The result at which we have arrived is not to be classed with the arbitrary methods which have been again and again propounded

to the famous Kelly Criterion [24]², although Kelly did not explicitly treat the St Petersburg game, and tangentially to Itô's lemma [22]. It appears as an exercise in a well-known text on information theory [14, Ex. 6.17]. Mainstream economics has ignored all this. A full and rigorous resolution of the paradox, including the epistemological significance of the shift from ensemble to time averages, was published recently by one of the present authors [43].

5.2 The Copenhagen experiment

At the end of Sec. ?? we arrived at a mapping between expected utility theory and ergodicity economics. The fact that such a mapping is possible is fascinating: careful thinking leads to almost identical mathematical expressions, whether we operate in the conceptual universe of 1738 or in that of today. Since the difference is first and foremost conceptual, it is fair to ask whether ergodicity economics is just another way of looking at the subject without fundamentally different or new results. That alone would be a wonderful achievement, as Rota put it when speaking about exterior algebra [?, p.48] [indiscrete thoughts]: “[It] is not meant to prove old facts, it is meant to disclose a new world. Disclosing new worlds is as worthwhile a mathematical enterprise as proving old conjectures.”

But we will see that the new world disclosed by ergodicity economics does more than prove old conjectures. Our first expeditions into this new world have revealed rich new lands and call for further ships to be sent out to sea.

So: different concepts, very nice. But do the differences between concepts allow an experiment, at least in principle, that would tell us something radically new. Specifically, since we can map ergodicity economics and expected utility theory under certain conditions, can an experiment be designed that has discriminating power between the two mapped models? If we mess with the conditions that allow the mapping, so that the two models are no longer equivalent, which model will be better at describing observed behaviour?

Without our knowledge, a group of neuroscientists at the Danish Research Center for Magnetic Resonance in Copenhagen was asking precisely this question.

They followed very closely the discussion in [48], where we had worked out in detail the correspondences between linear utility and additive dynamics; and between logarithmic utility and multiplicative dynamics. These correspondences provide the basis for the experiment: what if the dynamics of wealth could be controlled? With two artificial environments – one additive, one multiplicative – do people adjust their behaviour to be growth-optimal in each? That's the basic idea: observe the choices that people make under different wealth dynamics, and

to evade the difficulty of the Petersburg problem.... Formulae have often been proposed, which have possessed the one virtue of presenting a finite result... but they have often had no intelligible basis to rest upon, or... sufficient care has not been taken to draw a distinguishing line between the significance of the result obtained, and the different result arrived at when the mathematical expectation is calculated.” Sadly he chose to place these revolutionary remarks in an appendix of a college probability textbook.

²Kelly was similarly unimpressed with the mainstream and noted in his treatment of decision theory, which he developed from the perspective of information theory and which is identical to ergodicity economics with multiplicative dynamics, that the utility function is “too general to shed any light on the specific problems of communication theory.”

infer from those choices the ergodic growth rate (or utility function) that is being optimised.

A positive result – people changing their utility functions in response to the dynamics – would falsify expected utility theory (insofar as experiments falsify models). In the worldview of expected utility theory, the utility function encodes a personality. Supposedly, it’s a psychological or even neurological property of a human being, and not something that can change over the course of a few hours (the duration of the experiment).

A negative result – people not changing utility functions – would get ergodicity economics into trouble (at least in this experiment). In the worldview of ergodicity economics, the term “utility function” is a poorly chosen label for what’s really the ergodicity transformation that sits in the functional form of the relevant growth rate. This transformation is different for these two dynamics, as we’ve seen. And the postulate of ergodicity economics is that people behave so as to optimise the rate of change of the ergodicity transformation (*i.e.* the ergodic growth rate). If they keep using the same transformation, irrespective of dynamics, then this postulate would not describe real behaviour.

As we discuss the experiment below, we have to keep its limitations in mind. Of course, we could set up a flawed experiment. We could have a bug in the code to analyse the observations and accidentally introduce spurious differences between the inferred utility functions. Or we could set up an experiment that doesn’t achieve statistical significance to fish out of the noise the relevant differences in the ergodicity mappings that people optimise. Such experiments would not falsify anything; they would just be invalid.

The experiment is described in detail in [34]. Here we will only outline the setup. In the additive environment the subject was given a starting wealth of about \$150 and then made 312 choices of the (additive) form “would you rather toss a coin for winning \$40 or losing \$30; or for winning \$30 or losing \$20?”

In the multiplicative environment, the same subject was also given about \$150, and then made 312 choices of the (multiplicative) form “would you rather toss a coin for a 100% gain in wealth or a 70% loss; or for a 30% gain or a 20% loss?”

The choices were consequential: a single decision could lead to winning or losing several hundred real dollars.

We won’t go into the details of the data analysis, which is rather sophisticated in this case. Instead, we will only explain the logic behind what was done with the observations. Let’s work through a single observation – a single choice made by a subject. Let’s say the wealth of the subject is \$200). Imagine we observe the subject rejecting a proposed 50/50 gamble of winning \$100 or losing \$70 (for terminal wealth \$300 or \$130), and instead accepts a gamble of winning \$20 or losing \$10 (for terminal wealth \$220 or \$190). What does this tell us about the subject’s beliefs about the ergodicity transformation?

We’re testing ergodicity economics and expected utility theory. In both theories, the decision criterion is the expectation value of a change in some function of wealth. So the question is really: for what kind of $v(x)$ would the subject make the observed choice, assuming that the time-average ergodic growth rate is being optimised?

The time-average ergodic growth rate (or rate of change of expected utility)

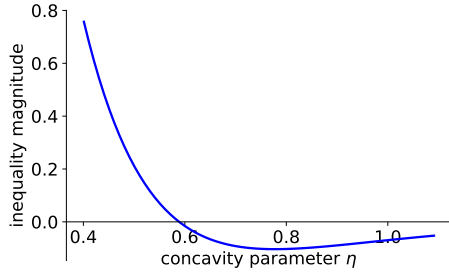


Figure 5.3: LHS of (Eq. 5.7). The inequality in (Eq. 5.7) is satisfied for any $\eta > 0.589$, wherefore the choice that the subject made indicates that an ergodicity transformation with $\eta > .589$ is being optimised. For instance, this observation is consistent with a logarithmic ergodicity transformation ($\eta = 1$).

{SC04}

{fig:eta_constraint}

is

$$g = \frac{\langle \Delta v(x) \rangle}{\Delta t}. \quad (5.4) \quad \{\text{eq:CPH_criterion}\}$$

We have to restrict the space of $v(x)$ in order to see roughly where in that space the subject sits. The restriction has to allow the two cases predicted by ergodicity economics: linear ergodicity transformation (utility function) for additive dynamics and logarithmic for multiplicative. The Copenhagen team chose to restrict $v(x)$ to a family of functions known in economics as auto-elastic utility functions, and in physics as q-logarithms, namely to

$$v(x; \eta) = \frac{x^{1-\eta} - 1}{1 - \eta}. \quad (5.5) \quad \{\text{eq:auto}\}$$

The parameter η controls the concavity of $v(x; \eta)$, with the special value $\eta = 0$ corresponding to linear $v(x)$, and $\eta = 1$ corresponding to logarithmic $v(x)$.

With the restriction to this family of ergodicity mappings, we can phrase the decision made by the subject in terms of beliefs about η . For instance, in the example choice we have

$$\frac{\frac{1}{2}[v(\$300) + v(\$130)] - v(\$200)}{\Delta t} < \frac{\frac{1}{2}[v(\$220) + v(\$190)] - v(\$200)}{\Delta t}. \quad (5.6)$$

To check when this inequality is satisfied, we subtract the RHS and simplify, to arrive at the condition

$$v(\$300) + v(\$130) - [v(\$220) + v(\$190)] < 0. \quad (5.7) \quad \{\text{eq:magnitude}\}$$

Figure 5.3 shows the LHS of (Eq. 5.7), plotted as a function of η . We see that for any $\eta > 0.589$ the inequality is satisfied. Assuming that the subject is maximising (Eq. 5.4), we conclude from this one observation that the subject believes $\eta > 0.589$. For example, the subject may believe that the dynamics are multiplicative ($\eta = 1$).

Each subject made 312 observed gambling choices. Each choice produces a similar statement – some constraint on η . Perhaps the next observation would show that the subject believes $\eta < 1.2$, which would bracket the subject's belief about the value of η into $0.589 < \eta < 1.2$.

Like we said, the fitting techniques that were used to analyse the data from the experiment are quite sophisticated. Where inconsistencies in a subject's behaviours are observed, *e.g.* of the type $\eta > 0.589$ and $\eta < 0.5$, this is interpreted as inaccuracy in the subject's estimate for η , in effect an inability of the subject to distinguish dynamics at this level of accuracy. The end result, after much

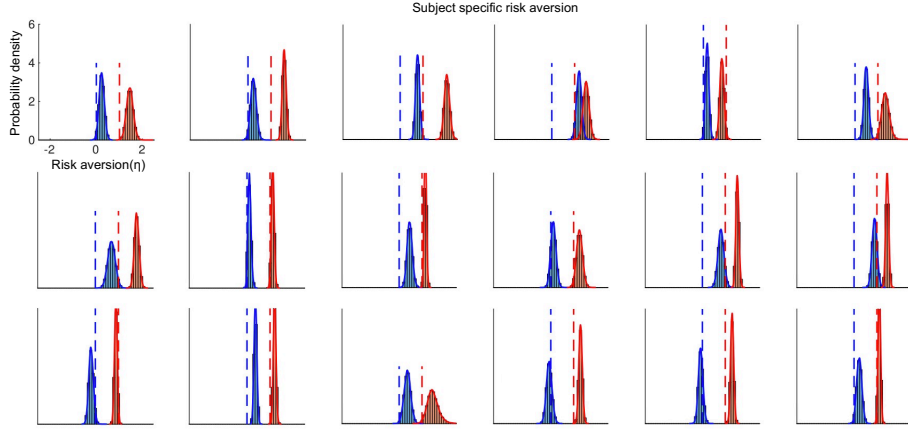


Figure 5.4: Posterior probability density functions for the parameter η in the Copenhagen experiment. Each panel is one individual, blue for the additive environment, red for multiplicative. Dashed lines are the null-model predictions of 0 and 1. All tested individuals changed behaviour noticeably in response to the wealth dynamic, in all cases the multiplicative environment led to the shift to the right predicted by ergodicity economics. Not all individuals are the same, but an overall pattern is clearly seen.

{fig:HulmeETAL}

computing, is a Bayesian posterior probability density function specifying the likelihood of a subject’s beliefs about η . Roughly, for each person in each environment this tells us how likely it is that the subject was optimizing expected changes in (Eq. 5.5) with different values for η , Fig. 5.4.

For details, see [34]. The statistical techniques used in the experiment are described in [27].

Expected-utility theory predicts that people are insensitive to changes in the dynamics. People may have wildly different utility functions, which would be reflected in wildly different best-fit values of η , but the dynamic setting should make no difference.

Ergodicity economics predicts something quite different. First of all, it predicts that the dynamic setting significantly changes the best-fit “utility function,” which is really the ergodicity transformation in the relevant ergodic growth rate. The effective utility function will be different for one and the same individual under additive dynamics and under multiplicative dynamics.

According to ergodicity economics, the direction of the change should go towards greater “risk aversion” for multiplicative dynamics – the ergodicity mapping is more concave there. The magnitude of the change in η should be about 1. And finally, if we take seriously the absolute null models of additive and multiplicative dynamics, the distributions should be centred near 0 for the additive setting and near 1 for the multiplicative setting.

Given the limitations of the experiment – for instance people only had a one-hour training phase to get used to a given dynamic environment – these predictions don’t look so bad. Of course the 11,232 individual choices summarized in Fig. 5.4 may be happenstance, or the experiment may be flawed in a way we don’t yet understand. So we might put it this way: saying that

expected-utility theory is based on good concepts corresponds to the belief that the difference between the red and blue curves is spurious.

5.3 Insurance

{section:Insurance}

The insurance contract is an important and ubiquitous type of economic transaction, which can be modelled as a gamble. Its significance is two-fold: insurance contracts transfer risk from one party to another, and this – the transfer, or sharing, of risk – is a key driver of the emergence of social structure at all levels of evolution Sec. 7.1. More mundanely, insurance is an important thing to understand because the selling and buying of insurance contracts (derivatives) constitutes the vast majority of global trade. Perhaps – a tantalising hypothesis – these two facts are not unrelated.

However, classically insurance poses a puzzle [45]. In the expected-wealth paradigm, insurance contracts shouldn't exist, because buying insurance would only be rational at a price at which it would be irrational to sell. More specifically:

1. To be viable, an insurer must charge an insurance premium of at least the expectation value of any claims that may be made against it, called the “net premium” [23, p. 1].
2. The insurance buyer therefore has to be willing to pay more than the expectation value of the claims he will receive (the net premium), so that an insurance contract may be successfully signed.
3. Under the expected-wealth paradigm it is irrational to pay more than the net premium, and therefore insurance contracts should not exist.

In this picture, an insurance contract can only ever be beneficial to one party. It has the anti-symmetric property that the expectation value of one party's gain is the expectation value of the other party's loss.

The puzzle is that insurance contracts are observed to exist. Why? Classical resolutions appeal to utility theory (*i.e.* psychology) and asymmetric information (*i.e.* deception). However, ergodicity economics naturally predicts contracts with a range of prices that increase the time-average growth rate for both buyer and seller. This is rather an important point. It means that, at a good price, insurance contracts are mutually beneficial – a theme we will return to in Sec. 7.1. We will also see that these voluntary contracts lead to a flow of cash from poor to rich – another theme we will return to, in Sec. ??.

We illustrate how this all works with an example drawn from maritime trade, in which the use of insurance has a very long history.³ A similar example was used by Bernoulli [6], incidentally in the paper containing his solution to the St. Petersburg paradox.

³Contracts between Babylonian traders and lenders were recorded around 1750 BC in the Code of Hammurabi. Chinese traders practised diversification by spreading cargoes across multiple vessels even earlier than this, in the third millennium BC.

5.4 Setup: a shipping insurance contract

We imagine a shipowner sending cargo from St Petersburg to Amsterdam, with the following parameters:

- owner's wealth, $x_{\text{own}} = \$100,000$;
- gain on safe arrival of cargo, $G = \$4,000$;
- probability ship will be lost, $p = 0.05$;
- replacement cost of the ship, $C = \$30,000$; and
- voyage time, $\delta t = 1$ month.

An insurer with wealth $x_{\text{ins}} = \$1,000,000$ proposes to insure the voyage for a fee, $F = \$1,800$. If the ship is lost, the insurer pays the owner $L = G + C$ to make him good on the loss of his ship and the profit he would have made.

We phrase the decision the owner is facing as a choice between two gambles.

Definition The owner's gambles

Sending the ship uninsured corresponds to gamble o1

$$q_1^{(\text{o1})} = G, \quad p_1^{(\text{o1})} = 1 - p; \quad (5.8)$$

$$q_2^{(\text{o1})} = -C, \quad p_2^{(\text{o1})} = p. \quad (5.9)$$

Sending the ship fully insured corresponds to gamble o2

$$q_1^{(\text{o2})} = G - F, \quad p_1^{(\text{o2})} = 1. \quad (5.10)$$

This is a trivial “gamble” because all risk has been transferred to the insurer.

We also model the insurer's decision whether to offer the contract as a choice between two gambles

Definition The insurer's gambles

Not insuring the ship corresponds to gamble i1, which is the null gamble

$$q_1^{(\text{i1})} = 0, \quad p_1^{(\text{i1})} = 1. \quad (5.11)$$

Insuring the ship corresponds to gamble i2

$$q_1^{(\text{i2})} = +F, \quad p_1^{(\text{i2})} = 1 - p; \quad (5.12)$$

$$q_2^{(\text{i2})} = -L + F, \quad p_2^{(\text{i2})} = p. \quad (5.13)$$

To establish the puzzle, we ask in the expected-wealth paradigm whether the owner should sign the contract, and whether the insurer should have proposed it.

Insurance in the expected-wealth paradigm

In the expected-wealth paradigm (corresponding to additive repetition under the time paradigm) decision makers maximise the rate of change of the expectation values of their wealths, according to (Eq. ??): Under this paradigm the owner collapses gamble o1 into the scalar

$$\bar{g}_a^{(o1)} = \frac{\langle \delta x \rangle}{\delta t} \quad (5.14)$$

$$= \frac{\langle q^{(o1)} \rangle}{\delta t} \quad (5.15)$$

$$= \frac{(1-p)G + p(-C)}{\delta t} \quad (5.16)$$

$$= \$2,300 \text{ per month}, \quad (5.17)$$

and gamble o2 into the scalar

$$\bar{g}_a^{o2} = \frac{\langle q^{(o2)} \rangle}{\delta t} \quad (5.18)$$

$$= \frac{(G - F)}{\delta t} \quad (5.19)$$

$$= \$2,200 \text{ per month}. \quad (5.20)$$

The difference, $\delta \bar{g}_a^o$, between the expected rates of change in wealth with and without a signed contract is the expected loss minus the fee per round trip,

$$\delta \bar{g}_a^o = \bar{g}_a^{o2} - \bar{g}_a^{o1} = \frac{pL - F}{\delta t}. \quad (5.21) \quad \{\text{eq:dri}\}$$

The sign of this difference indicates whether the insurance contract is beneficial to the owner. In the example this is not the case, $\delta \bar{g}_a^o = -\$100$ per month.

The insurer evaluates the gambles i1 and i2 similarly, with the result

$$\bar{g}_a^{(i1)} = \$0 \text{ per month}, \quad (5.22)$$

and

$$\bar{g}_a^{(i2)} = \frac{F - pL}{\delta t} \quad (5.23) \quad \{\text{eq:r}\}$$

$$= \$100 \text{ per month}. \quad (5.24)$$

Again we compute the difference – the net benefit to the insurer that arises from signing the contract

$$\delta \bar{g}_a^i = \bar{g}_a^{i2} - \bar{g}_a^{i1} = \frac{F - pL}{\delta t}. \quad (5.25) \quad \{\text{eq:dri}\}$$

In the example this is $\delta \bar{g}_a^i = \$100$ per month, meaning that in the world of the expected-wealth paradigm the insurer will offer the contract.

Because only one party (the insurer) is willing to sign, no contract will come into existence. We could think that we got the price wrong, and the contract would be signed if offered at a different fee. But this is not the case. Looking

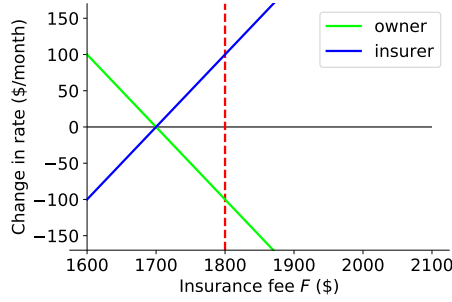


Figure 5.5: Change in the rate of change of expected wealth for the shipowner (green) and the insurer (blue) as a function of the insurance fee, F .

at (Eq. 5.21) and (Eq. 5.25) we notice the anti-symmetric relationship between the two expressions,

$$\delta \bar{g}_a^o = -\delta \bar{g}_a^i. \quad (5.26)$$

By symmetry, there can be no fee at which both expressions are positive. If one is positive, the other must be negative. The only point where this isn't the case is when both sides are exactly zero. Hence there are no circumstances in the world created by the expected-wealth paradigm under which both parties have an incentive to sign; there is merely one specific price at which neither party has a reason to refrain from signing. Insurance contracts cannot exist in this model world.

Fundamental insurance puzzle

In the model world created by expected-wealth maximisation, *no price exists* at which two parties will sign an insurance contract.

One party winning at the expense of the other makes insurance an unsavoury business in the expected-wealth paradigm. This is further illustrated in Fig. 5.5, which shows the change in the rate of change of expected wealth (the decision criterion) for both parties as a function of the fee, F . There is no price at which the decision variable is positive for both parties. The best they can do is to pick the price at which neither of them cares whether they sign or not.

In this picture, the existence of insurance contracts requires some asymmetry between the contracting parties, such as:

- different psychological attitudes to bearing risk;
- different access to information about the voyage;
- different assessments of the riskiness of the voyage;
- one party to deceive, coerce, or gull the other into a bad decision.

It is difficult to believe that this is truly the basis for a market of the size and global reach of the insurance market.

5.4.1 Ergodicity economics solution

Ergodicity economics resolves the insurance puzzle under many different types of dynamics – so long as the ergodicity mapping is concave. For simplicity we will

assume multiplicative dynamics, pausing to reflect on the difference to additive dynamics (which is equivalent to the expected-wealth paradigm that created the insurance puzzle). Assuming a multiplicative wealth dynamic models the ship owner as sending out a ship and a cargo whose values are proportional to his wealth at the start of each voyage. A rich owner who has had many successful voyages will send out more cargo, a larger ship, or perhaps a *flotilla*, while an owner to whom the sea has been a cruel mistress will send out a small vessel until his luck changes. Modelling wealth with an additive dynamic, the ship owner would send out the same amount of cargo on each journey, irrespective of his wealth. Shipping companies of the size of Evergreen or Maersk would be inconceivable under additive dynamics, where returns on successful investments are not reinvested.

Under multiplicative dynamics, the two parties seek to maximise

$$\bar{g}_m = \lim_{\Delta t \rightarrow \infty} \frac{\Delta v(x)}{\Delta t} = \frac{\langle \delta \ln x \rangle}{\delta t}, \quad (5.27)$$

where we have used the ergodic property of $\Delta v(x) = \Delta \ln x$.

The owner's time-average growth rate without insurance is

$$\bar{g}_m^{o1} = \frac{(1-p) \ln(x_{\text{own}} + G) + p \ln(x_{\text{own}} - C) - \ln(x_{\text{own}})}{\delta t} \quad (5.28)$$

or 1.9% per month. His time-average growth rate with insurance is

$$\bar{g}_m^{o2} = \frac{\ln(x_{\text{own}} + G - F) - \ln(x_{\text{own}})}{\delta t} \quad (5.29)$$

or 2.2% per month. This gives a net benefit for the owner of

$$\delta \bar{g}_m^o = \bar{g}_m^{o2} - \bar{g}_m^{o1} \approx +0.24\% \text{ per month.} \quad (5.30)$$

The time paradigm thus creates a world where the owner will sign the contract.

What about the insurer? Without insurance, the insurer plays the null gamble, and

$$\bar{g}_m^{i1} = \frac{0}{\delta t} \quad (5.31)$$

or 0% per month. His time-average growth rate with insurance is

$$\bar{g}_m^{i2} = \frac{(1-p) \ln(x_{\text{ins}} + F) + p \ln(x_{\text{ins}} + F - L) - \ln(x_{\text{ins}})}{\delta t} \quad (5.32)$$

or 0.0071% per month. The net benefit to the insurer is therefore also

$$\delta \bar{g}_m^i = \bar{g}_m^{i2} - \bar{g}_m^{i1} \quad (5.33)$$

i.e. 0.0071% per month. Unlike the expected wealth paradigm, ergodicity economics with multiplicative repetition creates a world where an insurance contract can exist – there exists a range of fees F at which both parties gain from signing the contract!

We view this as the fundamental resolution of the insurance puzzle.

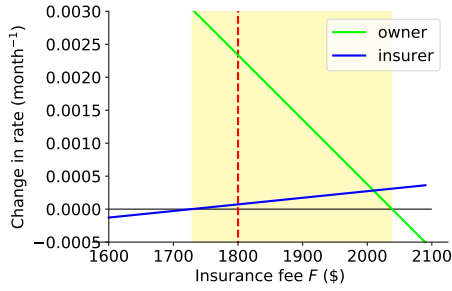


Figure 5.6: Change in the time-average growth rate of wealth for the shipowner (green) and the insurer (blue) as a function of the insurance fee, F . The mutually beneficial fee range is marked by the beige background.

{SC06}

{fig:ins_log}

Fundamental resolution of the insurance puzzle:

An insurance contract will be signed by both parties if doing so increases the time-average ergodic growth rate of both the seller's wealth and the buyer's wealth.

It requires no appeal to arbitrary utility functions or asymmetric circumstances, rather it arises naturally from the model of human decision-making that we have set out. Fig. 5.6 shows the mutually beneficial range of insurance fees predicted by our model.

Generalizing, the message of ergodicity economics is that business happens when both parties gain. In the world created by this model any agreement, any contract, any commercial interaction comes into existence because it is mutually beneficial.

5.4.2 The classical solution of the insurance puzzle

The classical solution of the insurance puzzle is identical to the classical solution of the St Petersburg paradox. Wealth is replaced by a non-linear utility function of wealth, which breaks the symmetry of the expected-wealth paradigm. While it is always true that $\delta \langle r \rangle_{\text{own}} = -\delta \langle r \rangle_{\text{ins}}$, the expected growth rates of non-linear utility functions don't share this anti-symmetry. A difference in the decision makers' wealths is sufficient, though often different utility functions are assumed for owner and insurer, which is a model that can create pretty much any behaviour. The downside of a model with this ability is, of course, that it makes no predictions – nothing is ruled out, so the model cannot be falsified. Popper would mark it as not science.

{section:The classical so

Part III

Macroeconomics

Chapter 6

People

{chapter:People}

The previous chapter developed a model of individual behaviour based on an assumed dynamic imposed on wealth. If we know the stochastic process that describes individual wealth, then we also know what happens at population level – each individual is represented by a realisation of the process, and we can compute the dynamics of wealth distributions. We answer questions about inequality and poverty in a model economy which, for the moment, contains no interactions between individuals. We also gain an understanding of when and how results for finite populations differ from those for the infinite ensemble.

6.1 Every man for himself

{section:Every_man}

We have seen that risk aversion constitutes optimal behaviour under the assumption of multiplicative wealth growth and over time scales that are long enough for systematic trends to be significant. In this chapter we will continue to explore our null model, **GBM**. By “explore” we mean that we will let the model generate its world – if individual wealth was to follow **GBM**, what kind of features of an economy would emerge? We will see that cooperation and the formation of social structure also constitute optimal behaviour.

GBM is more than a random variable. It’s a stochastic process, either a set of trajectories or a family of time-dependent random variables, depending on how we prefer to look at it. Both perspectives are informative in the context of economic modelling: from the set of trajectories we can judge what is likely to happen to an individual, *e.g.* by following a single trajectory for a long time; while the **PDF** of the random variable $x(t^*)$ at some fixed value of t^* tells us how wealth is distributed in our model.

We use the term wealth distribution to refer to the density function, $\mathcal{P}_x(x)$, and not the process of distributing wealth among people. This can be interpreted as follows. Imagine a population of N individuals. If I select a random individual, each having uniform probability $\frac{1}{N}$, then the probability of the selected individual having wealth greater than x is given by the CDF, $F_x(x) = \int_x^\infty \mathcal{P}_x(s)ds$. If N is large, then $\Delta x \mathcal{P}_x(x)N$ is the approximate number of individuals who have wealth between x and $x + \Delta x$. Thus, a broad wealth distribution with heavy tails indicates greater wealth inequality.

Examples:

- Under perfect equality everyone would have the same, meaning that the wealth distribution would be a Dirac delta function centred at the sample mean of x , that is

$$\mathcal{P}_x(x) = \delta(x - \langle x \rangle_N); \quad (6.1)$$

- Maximum inequality would mean that one individual owns everything and everyone else owns nothing, that is

$$\mathcal{P}_x(x) = \frac{N-1}{N} \delta(x-0) + \frac{1}{N} \delta(x-N \langle x \rangle_N). \quad (6.2)$$

6.1.1 Log-normal distribution

{section:Log-normal_wealth}

At a given time, t , **GBM** produces a random variable, $x(t)$, with a log-normal distribution whose parameters depend on t . (A log-normally distributed random variable is one whose logarithm is a normally distributed random variable.) If each individual’s wealth follows **GBM**,

$$dx = x(\mu dt + \sigma dW), \quad (6.3) \quad \{\text{eq:GBM}\}$$

with solution

$$x(t) = x(0) \exp \left[\left(\mu - \frac{\sigma^2}{2} \right) t + \sigma W(t) \right], \quad (6.4) \quad \{\text{eq:GBM_sol}\}$$

then we will observe a log-normal distribution of wealth at each moment in time:

$$\ln x(t) \sim \mathcal{N} \left(\ln x(0) + \left(\mu - \frac{\sigma^2}{2} \right) t, \sigma^2 t \right). \quad (6.5) \quad \{\text{eq:log-normal}\}$$

It will be convenient hereafter to assume the initial condition $x(0) = 1$ (and, therefore, $\ln x(0) = 0$) unless otherwise stated.

Note that the variance of $\ln x(t)$ increases linearly in time. We will develop an understanding of this shortly. As we will see, it indicates that any meaningful measure of inequality will grow over time in our simple model. To see what kind of a wealth distribution (Eq. 6.5) is, it is worth spelling out the log-normal PDF:

$$\mathcal{P}_x(x) = \frac{1}{x\sqrt{2\pi\sigma^2 t}} \exp\left(-\frac{[\ln x - (\mu - \frac{\sigma^2}{2})t]^2}{2\sigma^2 t}\right). \quad (6.6) \quad \{\text{eq:PDFx}\}$$

This distribution is the subject of a wonderful book [2], sadly out-of-print now. It will be useful to know some of its basic properties. Of particular importance is the expected wealth under this distribution. This is

$$\langle x(t) \rangle = \exp(\mu t) \quad (6.7) \quad \{\text{eq:exp_x}\}$$

or, equivalently, $\ln \langle x(t) \rangle = \mu t$. We could confirm this result by calculating $\langle x(t) \rangle = \int_0^\infty s \mathcal{P}_x(s) ds$, but this would be laborious. Instead we use a neat trick, courtesy of [?, Chapter 4.2], which will come in handy again in Sec. 7.2.6. The idea is to derive from the stochastic differential equation for x , like (Eq. 6.3), a solvable ordinary differential equation in the k^{th} moment, $\langle x^k \rangle$. For the first moment we do this simply by taking expectations of both sides of (Eq. 6.3). The noise term vanishes to turn the SDE for x into an ODE for $\langle x \rangle$:

$$\langle dx \rangle = \langle x(\mu dt + \sigma dW) \rangle \quad (6.8)$$

$$d\langle x \rangle = \langle x \rangle \mu dt + \sigma \overbrace{\langle dW \rangle}^{=0} \quad (6.9)$$

$$= \langle x \rangle \mu dt. \quad (6.10)$$

This is a very simple first-order linear differential equation, whose solution with initial condition $x(0) = 1$ is given by (Eq. 6.7).

For $\mu > 0$ the expected wealth grows exponentially over time, as do its median and variance:

$$\text{median}[x(t)] = \exp[(\mu - \sigma^2/2)t]; \quad (6.11) \quad \{\text{eq:median_x}\}$$

$$\text{var}[x(t)] = \exp(2\mu t)[\exp(\sigma^2 t) - 1]. \quad (6.12) \quad \{\text{eq:var_x}\}$$

6.1.2 Two growth rates

{section:two_rates}

We will recap briefly here one of our key ideas, covered in detail in Sec. 2.11, that the ensemble average of all possible trajectories of GBM grows at a different (faster) rate from that achieved by a single trajectory almost surely in the long-time limit. Understanding this difference was the key to developing a coherent theory of individual decision-making. We will see here that it is also crucial in understanding how wealth becomes distributed in a population of individuals whose wealths follow (Eq. 6.3) and, furthermore, how we can measure the inequality in such a distribution.

We recall from (Eq. 2.62) that the growth rate of the expected wealth is

$$g_{\langle \rangle} = \frac{d \ln \langle x \rangle}{dt} = \mu, \quad (6.13)$$

and from (Eq. 2.69) that the time-average growth rate of wealth is

$$\bar{g} = \frac{d\langle \ln x \rangle}{dt} = \mu - \frac{\sigma^2}{2}. \quad (6.14)$$

6.1.3 Measuring inequality

{section:Inequality_measu

In the case of GBM we have just seen how to compute the exact full wealth distribution, $\mathcal{P}_x(x)$. This is interesting but often we want only summary measures of the distribution. One summary measure of particular interest to economists is inequality. How much inequality is there in a distribution like (Eq. 6.5)? And how does this quantity increase over time under GBM, as we have suggested?

Clearly, to answer these questions, we must quantify “inequality”. In this section, and also in [1], we develop a natural way of measuring it, which makes use of the two growth rates we identified for the non-ergodic process. We will see that a particular inequality measure, known to economists as Theil’s second index of inequality [58], is the difference between typical wealth (growing at the time-average growth rate) and average wealth (growing at the ensemble-average growth rate) in our model. Thus, the difference between the time average and ensemble average, the essence of ergodicity breaking, is the fundamental driver of the dynamics of inequality.

The two limits of inequality are easily identified: minimum inequality means that everyone has the same wealth, and maximum inequality means that one individual has all the wealth and everyone else has nothing. (This assumes that wealth cannot become negative.) Quantifying inequality in any other distribution is reminiscent of the gamble problem. Recall that for gambles we wanted make statements of the type “this gamble is more desirable than that gamble”. We did this by collapsing a distribution to a scalar. Depending on the question that was being asked, the appropriate way of collapsing the distribution and the resulting scalar can be different (the scalar relevant to an insurance company may not be relevant to an individual). In the case of inequality we also have a distribution – the wealth distribution – and we want to make statements of the type “this distribution is more unequal than that distribution”. Again, this is done by collapsing the distribution to a scalar, and again many different choices of collapse and resulting scalar are possible. The Gini coefficient is a particularly well-known scalar of this type, the 80/20 ratio another. Indeed, there is a whole menagerie of inequality measures.

In this context the expectation value is an important quantity. For instance, if everyone has the same wealth, then everyone will own the average, $x_i = \langle x \rangle_N$, which converges to the expectation value for large N . Also, whatever the distribution of wealth, the total wealth is $N \langle x \rangle_N$ which converges to $N \langle x \rangle$ as N grows large. The growth rate of the expectation value, $g_{\langle x \rangle}$, thus tells us how fast the average wealth and the total population wealth grow with probability one in a large ensemble. The time-average growth rate, $\bar{g}$, on the other hand, tells us how fast an individual’s wealth grows almost surely in the long run. If the typical individual’s wealth grows more slowly than the expectation value of wealth, then there must be atypical individuals with very large wealths to account for the difference. This insight suggests the following measure of inequality.

Definition Inequality, J , is the quantity whose growth rate is the

difference between expectation-value and time-average growth rates,

$$\frac{dJ}{dt} = g_{\langle \rangle} - \bar{g}. \quad (6.15) \quad \{\text{eq:dJ}\}$$

Equation (6.15) defines the dynamic of inequality, and inequality itself is found by integrating over time:

$$J(t) = \int_0^t ds [g_{\langle \rangle}(s) - \bar{g}(s)]. \quad (6.16) \quad \{\text{eq:J}\}$$

This definition may be used for dynamics other than GBM, as we shall discuss in Sec. 6.1.6. Whatever the wealth dynamic, typical minus average growth rates are informative of the dynamic of inequality. Within the GBM framework we can write the difference in growth rates as

$$\frac{dJ}{dt} = \frac{d \ln \langle x \rangle}{dt} - \frac{d \langle \ln x \rangle}{dt} \quad (6.17) \quad \{\text{eq:J_dyn}\}$$

and integrate over time to get

$$J(t) = \ln \langle x(t) \rangle - \langle \ln x(t) \rangle. \quad (6.18) \quad \{\text{eq:J_x}\}$$

This quantity is known as the mean logarithmic deviation (MLD) or Theil's second index of inequality [58]. This is rather remarkable. Our general inequality measure, (Eq. 6.16), evaluated for the specific case of GBM, turns out to be a well-known measure of inequality that economists have identified independently, without considering non-ergodicity and ensemble average and time average growth rates. Merely by insisting on measuring inequality well, Theil used the GBM model implicitly.¹

Substituting the known values of the two growth rates into (Eq. 6.15) and integrating, we can evaluate the Theil inequality as a function of time:

$$J(t) = \frac{\sigma^2}{2} t. \quad (6.19) \quad \{\text{eq:J_t}\}$$

Thus we see that, in GBM, our measure of inequality increases indefinitely.

6.1.4 Wealth condensation

{section:condensation}

The log-normal distribution generated by GBM broadens indefinitely, (Eq. 6.12). Likewise, the inequality present in the distribution – measured as the time-integrated difference between ensemble and time average growth rates – grows without bound. A related property of GBM is the evolution towards wealth condensation. Wealth condensation means that a single individual will own a non-zero fraction of the total wealth in the population in the limit of large N , see *e.g.* [10]. In the present case an arbitrarily large share of the total wealth will be owned by an arbitrarily small share of the population.

One simple way of seeing this is to calculate the fraction of the population whose wealths are less than the expected wealth, *i.e.* $x(t) < \exp(\mu t)$. To do

¹Like Kelly's ideas about gambling [24], Theil's inequality measures were developed using information theory.

this, we define a new random variable, $z(t)$, whose distribution is the standard normal:

$$z(t) \equiv \frac{\ln x(t) - (\mu - \sigma^2/2)t}{\sigma t^{1/2}} \sim \mathcal{N}(0, 1). \quad (6.20)$$

We want to know the mass of the distribution with $\ln x(t) < \mu t$. This is equivalent to $z < \sigma t^{1/2}/2$, so the fraction below the expected wealth is

$$\Phi\left(\frac{\sigma t^{1/2}}{2}\right), \quad (6.21)$$

where Φ is the CDF of the standard normal distribution. This increases over time, tending to one as $t \rightarrow \infty$.

6.1.5 Rescaled wealth

{section:rescaled}

Economists have arrived at many inequality measures, and have drawn up a list of conditions that particularly useful measures of inequality satisfy. Such measures are called “relative measures” [54, Appendix 4] and J is one of them.

One of the conditions is that inequality measures should not change when x is divided by the same factor for everyone. Since we are primarily interested in inequality in this section, we can remove absolute wealth levels from the analysis and study an object called the rescaled wealth.

Definition The rescaled wealth,

$$y_i(t) = \frac{x_i(t)}{\langle x(t) \rangle_N}, \quad (6.22) \quad \{\text{eq:rescaled}\}$$

is the proportion of the sample mean wealth – *i.e.* the wealth averaged over the finite population – owned by an individual.

This quantity is useful because its numerical value does not depend on the currency used: it is a dimensionless number. Thus if my rescaled wealth is $y_i(t) = 1/2$, it means that my wealth is half the average wealth, irrespective of whether I measure it in Kazakhstani Tenge or in Swiss Francs. The sample mean rescaled wealth is easily calculated:

$$\langle y_i(t) \rangle_N = \left\langle \frac{x(t)}{\langle x(t) \rangle_N} \right\rangle_N = 1. \quad (6.23)$$

If the population size, N , is large enough, then we might expect the sample mean wealth, $\langle x(t) \rangle_N$, to be close to the ensemble average, $\langle x(t) \rangle$, which is simply its $N \rightarrow \infty$ limit. We will discuss more carefully when this approximation holds for wealths following GBM in Sec. 6.2. Let’s assume for now that it does. The rescaled wealth is then well approximated as

$$y_i(t) = \frac{x_i(t)}{\langle x(t) \rangle} = x_i(t) \exp(-\mu t). \quad (6.24)$$

Now that we have an expression for y in terms of x and t , we can derive the dynamic for rescaled wealth using Itô’s formula (just as we did to find the

wealth dynamic for a general utility function in Sec. ??). We start with

$$dy = \frac{\partial y}{\partial t} dt + \frac{\partial y}{\partial x} dx + \frac{1}{2} \frac{\partial^2 y}{\partial x^2} dx^2 \quad (6.25)$$

$$= -\mu y dt + \frac{y}{x} dx, \quad (6.26) \quad \{\text{eq:ysde}\}$$

and then substitute (Eq. 6.3) for dx to get

$$dy = y\sigma dW. \quad (6.27) \quad \{\text{eq:GBM}_y\}$$

So $y(t)$ follows a very simple **GBM** with zero drift and volatility σ . This means that rescaled wealth, like wealth, has an ever-broadening log-normal distribution:

$$\ln y(t) \sim \mathcal{N}\left(-\frac{\sigma^2}{2}t, \sigma^2 t\right). \quad (6.28) \quad \{\text{eq:log-normal}_y\}$$

Finally, noting that $\langle \ln y \rangle = \langle \ln x \rangle - \ln \langle x \rangle$ gives us a simple expression for our inequality measure in (Eq. 6.18) in terms of the rescaled wealth:

$$J(t) = -\langle \ln y \rangle. \quad (6.29)$$

6.1.6 u -normal distributions and Jensen's inequality

`{section:jensen}`

So far we have confined our analysis to **GBM**, where wealths follow the dynamic specific by (Eq. 6.3). However, as we discussed in the context of gambles, other wealth dynamics are possible. In particular, we explored the dynamics corresponding to invertible utility functions, where utility executes a **BM** with drift as in (Eq. ??):

$$du = a_u dt + b_u dW. \quad (6.30)$$

Under this dynamic, utility is normally distributed,

$$u(x(t)) \sim \mathcal{N}(a_u t, b_u^2 t), \quad (6.31)$$

and we can say that wealth has a “ u -normal” distribution. For **GBM**, the corresponding utility function is, as we know, $u(x) = \ln x$, and u -normal becomes log-normal.

Replacing the logarithm in (Eq. 6.18) by the general utility function gives a general expression for our wealth inequality measure,

$$J_u(t) = u(\langle x(t) \rangle) - \langle u(x(t)) \rangle. \quad (6.32) \quad \{\text{eq:J}_u\}$$

It's not easy to write a general expression for $J_u(t)$ in only t and the model parameters a_u and b_u , because it would involve the solution of the general wealth dynamic in (Eq. 4.60). However, we can still say something about how inequality evolves. Let's see what happens if we start with perfect equality, $x_i(0) = x_0$ with x_0 fixed, and then let wealths evolve a little to $x(\Delta t) = x_0 + \Delta x$, where Δx is a random wealth increment generated by the wealth dynamic. The change in inequality would be

$$\Delta J_u = u(\langle x_0 + \Delta x \rangle) - \langle u(x_0 + \Delta x) \rangle, \quad (6.33)$$

since $u(\langle x_0 \rangle) = \langle u(x_0) \rangle = u(x_0)$.

We can now appeal to Jensen's inequality: if u is a concave function, like the logarithm, then $\Delta J_u \geq 0$; while if u is convex, like the curious exponential in (Eq. 4.86), then $\Delta J_u \leq 0$. The only cases for which $\Delta J_u = 0$ are if Δx is non-random or if u is linear. Thus, it is both randomness and the nonlinearity of the utility function – or ergodicity transformation, if we make that equivalence – that creates a difference in growth rates and generates inequality.

6.1.7 power-law resemblance

{section:power_law}

It is an established empirical observation [39] that the upper tails of real wealth distributions look more like a power-law than a log-normal. Our trivial model does not strictly reproduce this feature, but it is instructive to compare the log-normal distribution to a power-law distribution. A power-law PDF has the asymptotic form

$$\mathcal{P}_x(x) = x^{-\alpha}, \quad (6.34) \quad \{\text{eq:power_law}\}$$

for large arguments x . This implies that the logarithm of the PDF is proportional to the logarithm of its argument, $\ln \mathcal{P}_x(x) = -\alpha \ln x$. Plotting one against the other will yield a straight line, the slope being the exponent $-\alpha$.

Determining whether an empirical observation is consistent with such behaviour is difficult because the behaviour to be observed is in the tail (large x) where data are, by definition, sparse. A quick-and-dirty way of checking for possible power-law behaviour is to plot an empirical PDF against its argument on log-log scales, look for a straight line, and measure the slope. However, plotting any distribution on any type of scales results in some line. It may not be a straight line but it will have some slope everywhere. For a known distribution (power-law or not) we can interpret this slope as a local apparent power-law exponent.

What is the local apparent power-law exponent of a log-normal wealth distribution near the expectation value $\langle x \rangle = \exp(\mu t)$, *i.e.* in the upper tail where approximate power-law behaviour has been observed empirically? The logarithm of (Eq. 6.6) is

$$\ln \mathcal{P}(x) = -\ln \left(x \sqrt{2\pi\sigma^2 t} \right) - \frac{[\ln x - (\mu - \frac{\sigma^2}{2})t]^2}{2\sigma^2 t} \quad (6.35)$$

$$= -\ln x - \frac{\ln(2\pi\sigma^2 t)}{2} - \frac{(\ln x)^2 - 2(\mu - \frac{\sigma^2}{2})t \ln x + (\mu - \frac{\sigma^2}{2})^2 t^2}{2\sigma^2 t}. \quad (6.36)$$

Collecting terms in powers of $\ln x$ we find

$$\ln \mathcal{P}(x) = -\frac{(\ln x)^2}{2\sigma^2 t} + \left(\frac{\mu}{\sigma^2} - \frac{3}{2} \right) \ln x - \frac{\ln(2\pi\sigma^2 t)}{2} - \frac{(\mu - \frac{\sigma^2}{2})^2 t}{2\sigma^2} \quad (6.37)$$

with local slope, *i.e.* apparent exponent,

$$\frac{d \ln \mathcal{P}(x)}{d \ln x} = -\frac{\ln x}{\sigma^2 t} + \frac{\mu}{\sigma^2} - \frac{3}{2}. \quad (6.38)$$

Near $\langle x \rangle$, $\ln x \sim \mu t$ so that the first two terms cancel approximately. Here the distribution will resemble a power-law with exponent $-3/2$ when plotted on doubly logarithmic scales. (The distribution will also look like a power-law where the first term is much smaller than the others, *e.g.* where $\ln x \ll \sigma^2 t$.)

We don't believe that such empirically observed power-laws are merely a manifestation of this mathematical feature. Important real-world mechanisms that broaden real wealth distributions, *i.e.* concentrate wealth, are missing from the null model. However, it is interesting that the trivial model of **GBM** reproduces so many qualitative features of empirical observations.

6.2 Finite populations

{section:finite_population}

So far we have considered the properties of the random variable, $x(t)$, generated by **GBM** at a fixed time, t . Most of the mathematical objects we have discussed are, strictly speaking, relevant only in the limit $N \rightarrow \infty$, where N is the number of realisations of this random variable. For example, the expected wealth, $\langle x(t) \rangle$, is the limit of the sample mean wealth

$$\langle x(t) \rangle_N \equiv \frac{1}{N} \sum_{i=1}^N x_i(t), \quad (6.39) \quad \{\text{eq:sample}\}$$

as the sample size, N , grows large. In reality, human populations can be very large, say $N \sim 10^7$ for a nation state, but they are most certainly finite. Therefore, we need to be diligent and ask what the effects of this finiteness are. In particular, we will focus on the sample mean wealth under **GBM**. For what values of μ , σ , t , and N is this well approximated by the expectation value? And when it is not, what does it resemble?

6.2.1 Sums of log-normals

{section:sketch}

In [49] we studied the sample mean of **GBM**, which we termed the “partial ensemble average” (PEA). This is the average of N independent realisations of the random variable $x(t)$, (Eq. 6.39). Here we sketch out some simple arguments about how this object depends on N and t .

Considering the two growth rates in Sec. 6.1.2, we anticipate the following tension:

- A) for large N , the PEA should resemble the expectation value, $\exp(\mu t)$;
- B) for long t , all trajectories in the sample – and, therefore, the sample mean itself – should grow like $\exp[(\mu - \sigma^2/2)t]$.

Situation A – when a sample mean resembles the corresponding expectation value – is known in statistical physics as “self-averaging.” A simple strategy for estimating when this occurs is to look at the relative variance of the PEA,

$$R \equiv \frac{\text{var}[\langle x(t) \rangle_N]}{\langle \langle x(t) \rangle_N \rangle^2}. \quad (6.40) \quad \{\text{eq:rel_var}\}$$

To be explicit, here the $\langle \cdot \rangle$ and $\text{var}(\cdot)$ operators, without N as a subscript, refer to the mean and variance over all possible PEAs. The PEAs themselves, taken over finite samples of size N , are denoted $\langle \cdot \rangle_N$. Equation (6.40) simplifies to

$$R = \frac{\frac{1}{N} \text{var}[x(t)]}{\langle x(t) \rangle^2}, \quad (6.41)$$

into which we insert (Eq. 6.7) and (Eq. 6.12) to get an expression in terms of the GBM model parameters:

$$R = \frac{e^{\sigma^2 t} - 1}{N}. \quad (6.42) \quad \{\text{eq:rel_var_N}\}$$

If $R \ll 1$, then the PEA will likely be close to its own expectation value, which is equal to the expectation value of the GBM. Therefore, $\langle x(t) \rangle_N \approx \langle x(t) \rangle$ when

$$t < \frac{\ln N}{\sigma^2}. \quad (6.43) \quad \{\text{eq:short_t}\}$$

This hand-waving tells us roughly when the large-sample – or, as we see from (Eq. 6.43), short-time or low-volatility – self-averaging regime holds. A more careful estimate of the cross-over time in (Eq. 6.59) is a factor of 2 larger, as we shall see in Sec. 6.2.2, but the scaling is identical.

For $t > \ln N / \sigma^2$, the growth rate of the PEA transitions from μ to its $t \rightarrow \infty$ limit of $\mu - \sigma^2/2$ (Situation B). Another way of viewing this is to think about what dominates the average. For early times in the process, all trajectories are close together and none dominate the PEA. However, as time goes by the distribution broadens exponentially. Since each trajectory contributes with the same weight to the PEA, after some time the PEA will be dominated by the maximum in the sample,

$$\langle x(t) \rangle_N \approx \frac{1}{N} \max_{i=1}^N \{x_i(t)\}. \quad (6.44)$$

All of these features are illustrated in Fig. 6.1.

Self-averaging stops when even the “luckiest” trajectory is no longer close to the expectation value $\exp(\mu t)$. This is guaranteed to happen eventually because the probability for a trajectory to reach $\exp(\mu t)$ decreases towards zero as t grows, as we saw in Sec. 6.1.4. Naturally, this takes longer for larger samples, which have more chances to contain a lucky trajectory.

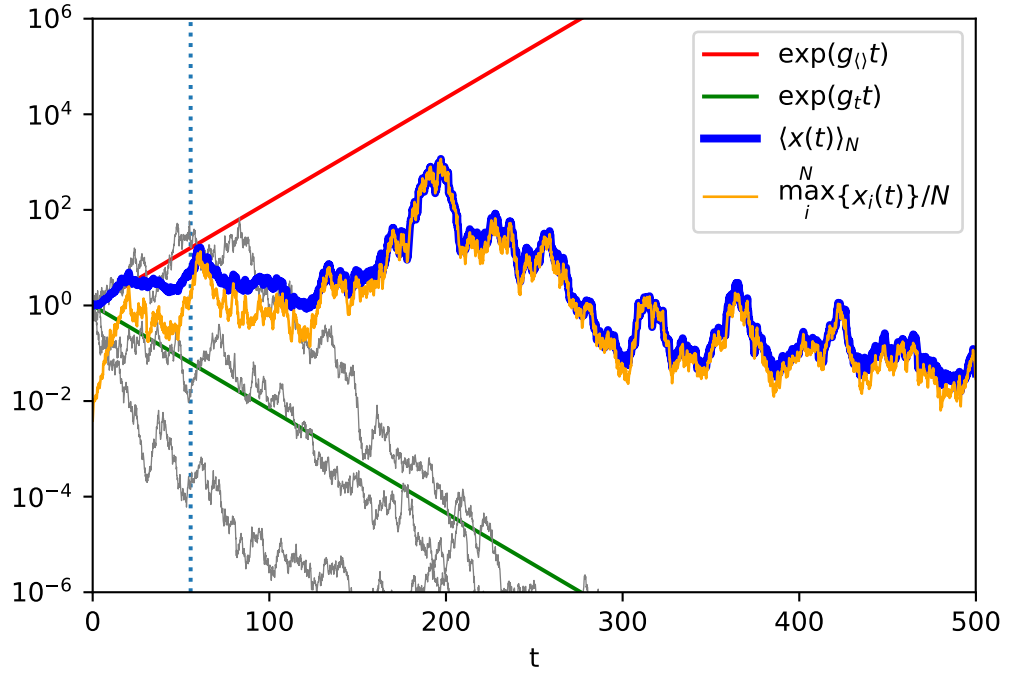


Figure 6.1: PEA and maximum in a finite ensemble of size $N = 256$. **Red line:** expectation value $\langle x(t) \rangle$. **Green line:** exponential growth at the time-average growth rate. In the $t \rightarrow \infty$ limit all trajectories grow at this rate. **Yellow line:** contribution of the maximum value of any trajectory at time t to the PEA. **Blue line:** PEA $\langle x(t) \rangle_N$. **Vertical line:** Crossover – for $t > t_c = \frac{2 \ln N}{\sigma^2}$ the maximum begins to dominate the PEA (the yellow line approaches the blue line). **Grey lines:** randomly chosen trajectories – any typical trajectory soon grows at the time-average growth rate. **Parameters:** $N = 256$, $\mu = 0.05$, $\sigma = \sqrt{0.2}$.

{fig:trajectories}

In [49] we analysed PEAs of GBM analytically and numerically. Using (Eq. 6.4) the PEA can be written as

$$\langle x \rangle_N = \frac{1}{N} \sum_{i=1}^N \exp \left[\left(\mu - \frac{\sigma^2}{2} \right) t + \sigma W_i(t) \right], \quad (6.45) \quad \{\text{eq:PEA}\}$$

where $\{W_i(t)\}_{i=1\dots N}$ are N independent realisations of the Wiener process. Taking the deterministic part out of the sum, we can write

$$\langle x \rangle_N = \exp \left[\left(\mu - \frac{\sigma^2}{2} \right) t \right] \frac{1}{N} \sum_{i=1}^N \exp \left(t^{1/2} \sigma \xi_i \right), \quad (6.46) \quad \{\text{eq:PEA}_2\}$$

where $\{\xi_i\}_{i=1\dots N}$ are N independent standard normal variates.

We found that typical trajectories of PEAs grow at $g_{\langle \rangle}$ up to a time t_c that is logarithmic in N , meaning $t_c \propto \ln N$. This is consistent with our analytical sketch. After this time, typical PEA trajectories begin to deviate from expectation-value behaviour, and eventually their growth rate converges to g_t . While the two limiting behaviours $N \rightarrow \infty$ and $t \rightarrow \infty$ can be computed exactly, what happens in between is less straightforward. The PEA is a random object outside these limits.

A quantity of crucial interest to us is the exponential growth rate experienced by the PEA,

$$g_{\text{est}}(t, N) \equiv \frac{\ln(\langle x(t) \rangle_N) - \ln(x(0))}{t - 0} = \frac{1}{t} \ln(\langle x(t) \rangle_N). \quad (6.47) \quad \{\text{eq:g_est}\}$$

In [49] we proved that the $t \rightarrow \infty$ limit for any (finite) N is the same as for the case $N = 1$,

$$\lim_{t \rightarrow \infty} g_{\text{est}}(t, N) = \mu - \frac{\sigma^2}{2} \quad (6.48) \quad \{\text{eq:g_est}_2\}$$

for all $N \geq 1$. Substituting (Eq. 6.46) in (Eq. 6.47) produces

$$g_{\text{est}}(t, N) = \mu - \frac{\sigma^2}{2} + \frac{1}{t} \ln \left(\frac{1}{N} \sum_{i=1}^N \exp(t^{1/2} \sigma \xi_i) \right) \quad (6.49)$$

$$= \mu - \frac{\sigma^2}{2} - \frac{\ln N}{t} + \frac{1}{t} \ln \left(\sum_{i=1}^N \exp(t^{1/2} \sigma \xi_i) \right). \quad (6.50) \quad \{\text{eq:g_est}_4\}$$

We didn't look in [49] at the expectation value of $g_{\text{est}}(t, N)$ for finite time and finite samples, but it's an interesting object that depends on N and t but is not stochastic. Note that this is not g_{est} of the expectation value, which would be the $N \rightarrow \infty$ limit of (Eq. 6.47). Instead it is the $S \rightarrow \infty$ limit,

$$\langle g_{\text{est}}(t, N) \rangle = \frac{1}{t} \langle \ln(\langle x(t) \rangle_N) \rangle = f(N, t), \quad (6.51) \quad \{\text{eq:g_est}_3\}$$

where, as previously, $\langle \cdot \rangle$ without subscript refers to the average over all possible samples, *i.e.* $\lim_{S \rightarrow \infty} \langle \cdot \rangle_S$. The last two terms in (Eq. 6.50) suggest an exponential relationship between ensemble size and time. The final term is a tricky stochastic object on which the properties of the expectation value in (Eq. 6.51) will hinge. This term will be the focus of our attention: the sum of exponentials of normal random variates or, equivalently, log-normal variates.

6.2.2 The random energy model

{section:REM}

Since the publication of [49] we have learned, thanks to discussions with J.-P. Bouchaud, that the key object in (Eq. 6.50) – the sum log-normal random variates – has been of interest to the mathematical physics community since the 1980s. The reason for this is Derrida’s random energy model [15, 16], which we will describe here. This section is very “physicsy” and the key results are in (Eq. 6.64) and (Eq. 6.65) – it’s safe to skip to them, if you like.

The model is defined as follows. Imagine a system whose energy levels are $2^K = N$ normally-distributed random numbers, ξ_i (corresponding to K spins). This is a very simple model of a disordered system, such as a spin glass, the idea being that the system is so complicated that we “give up” and simply model its energy levels as realisations of a random variable. (We denote the number of spins by K and the number of resulting energy levels by N , while Derrida uses N for the number of spins). The partition function is then

$$Z = \sum_{i=1}^N \exp \left(\beta J \sqrt{\frac{K}{2}} \xi_i \right), \quad (6.52) \quad \{\text{eq:Z}\}$$

where the inverse temperature, β , is measured in appropriate units, and the scaling in K is chosen so as to ensure an extensive thermodynamic limit [15, p. 79]. J is a constant that will be determined below. The logarithm of the partition function gives the Helmholtz free energy,

$$F = -\frac{\ln Z}{\beta} \quad (6.53)$$

$$= -\frac{1}{\beta} \ln \left[\sum_{i=1}^N \exp \left(\beta J \sqrt{\frac{K}{2}} \xi_i \right) \right]. \quad (6.54) \quad \{\text{eq:F}\}$$

Like the growth rate estimator in (Eq. 6.47), this involves a sum of log-normal variates and, indeed, we can rewrite (Eq. 6.50) as

$$g_{\text{est}} = \mu - \frac{\sigma^2}{2} - \frac{\ln N}{t} - \frac{\beta F}{t}, \quad (6.55) \quad \{\text{eq:ggest_5}\}$$

which is valid provided that

$$\beta J \sqrt{\frac{K}{2}} = \sigma t^{1/2}. \quad (6.56) \quad \{\text{eq:map}\}$$

Equation (6.56) does not give a unique mapping between the parameters of our GBM, (σ, t) , and the parameters of the REM, (β, K, J) . Equating (up to multiplication) the constant parameters, σ and J , in each model gives us a specific mapping:

$$\sigma = \frac{J}{\sqrt{2}} \quad \text{and} \quad t^{1/2} = \beta \sqrt{K}. \quad (6.57) \quad \{\text{eq:choice_1}\}$$

The expectation value of g_{est} is interesting. The only random object in (Eq. 6.55) is F . Knowing $\langle F \rangle$ thus amounts to knowing $\langle g_{\text{est}} \rangle$. In the statistical mechanics of the random energy model, $\langle F \rangle$ is of key interest and so much about it is known. We can use this knowledge thanks to the mapping between the two problems.

Derrida identifies a critical temperature,

$$\frac{1}{\beta_c} \equiv \frac{J}{2\sqrt{\ln 2}}, \quad (6.58) \quad \{\text{eq:beta_c}\}$$

above and below which the expected free energy scales differently with K and β . This maps to a critical time scale in GBM,

$$t_c = \frac{2 \ln N}{\sigma^2}, \quad (6.59) \quad \{\text{eq:t_c}\}$$

with high temperature ($1/\beta > 1/\beta_c$) corresponding to short time ($t < t_c$) and low temperature ($1/\beta < 1/\beta_c$) corresponding to long time ($t > t_c$). Note that t_c in (Eq. 6.59) scales identically with N and σ as the transition time, (Eq. 6.43), in our sketch.

In [15], $\langle F \rangle$ is computed in the high-temperature (short-time) regime as

$$\langle F \rangle = E - S/\beta \quad (6.60)$$

$$= -\frac{K}{\beta} \ln 2 - \frac{\beta K J^2}{4}, \quad (6.61) \quad \{\text{eq:F_2}\}$$

and in the low-temperatures (long-time) regime as

$$\langle F \rangle = -KJ\sqrt{\ln 2}. \quad (6.62) \quad \{\text{eq:F_3}\}$$

Short time

We look at the short-time behavior first (high $1/\beta$, (Eq. 6.61)). The relevant computation of the entropy S in [15] involves replacing the number of energy levels $N(E)$ by its expectation value $\langle N(E) \rangle$. This is justified because the standard deviation of this number is $\sqrt{N}$ and relatively small when $\langle N(E) \rangle > 1$, which is the interesting regime in Derrida's case.

For spin glasses, the expectation value of F is interesting, supposedly, because the system may be self-averaging and can be thought of as an ensemble of many smaller sub-systems that are essentially independent. The macroscopic behavior is then given by the expectation value.

Taking expectation values and substituting from (Eq. 6.61) in (Eq. 6.55) we find

$$\langle g_{\text{est}} \rangle^{\text{short}} = \mu - \frac{\sigma^2}{2} + \frac{1}{t} \frac{KJ^2}{4T^2}. \quad (6.63) \quad \{\text{eq:g_est_6}\}$$

From (Eq. 6.56) we know that $t = \frac{KJ^2}{2\sigma^2 T^2}$, which we substitute to find

$$\langle g_{\text{est}} \rangle^{\text{short}} = \mu. \quad (6.64) \quad \{\text{eq:g_est_7}\}$$

This is the correct behavior in the short-time regime.

Long time

Next, we turn to the expression for the long-time regime (low temperature, (Eq. 6.62)). Again taking expectation values and substituting, this time from (Eq. 6.62) in (Eq. 6.55), we find for long times

$$\langle g_{\text{est}} \rangle^{\text{long}} = \mu - \frac{\sigma^2}{2} - \frac{\ln N}{t} + \sqrt{\frac{2 \ln N}{t}} \sigma, \quad (6.65) \quad \{\text{eq:g_est_8}\}$$

which has the correct long-time asymptotic behavior. The form of the correction to the time-average growth rate in (Eq. 6.65) is consistent with [49] and [51], where it was found that approximately $N = \exp(t)$ systems are required for ensemble-average behavior to be observed for a time t , so that the parameter $\ln N/t$ controls which regime dominates. If the parameter is small, then (Eq. 6.65) indicates that the long-time regime is relevant.

Figure 6.2 is a direct comparison between the results derived here, based on [15], and numerical results using the same parameter values as in [49], namely $\mu = 0.05$, $\sigma = \sqrt{0.2}$, $N = 256$ and $S = 10^5$.

Notice that $\langle g_{\text{est}} \rangle$ is not the (local) time derivative $\frac{\partial}{\partial t} \langle \ln(\langle x \rangle_N) \rangle$, but a time-average growth rate, $\left\langle \frac{1}{t} \ln \left(\frac{\langle x(t) \rangle_N}{\langle x(0) \rangle_N} \right) \right\rangle$. It is remarkable that the expectation value $\langle g_{\text{est}}(N, t) \rangle$ so closely reflects the median, $q_{0.5}$, of $\langle x \rangle_N$, *i.e.*

$$q_{0.5}(\langle x(t) \rangle_N) \approx \exp(\langle g_{\text{est}}(N, t) \rangle t). \quad (6.66) \quad \{\text{eq:quant_ave}\}$$

In [48] it was discussed in detail that $g_{\text{est}}(1, t)$ is an ergodic observable for (Eq. 6.3), in the sense that $\langle g_{\text{est}}(1, t) \rangle = \lim_{t \rightarrow \infty} g_{\text{est}}$. The relationship in (Eq. 6.66) is far more subtle. The typical behavior of GBM PEAs is complicated outside the limits $N \rightarrow \infty$ or $t \rightarrow \infty$, where growth rates are time-dependent. This complicated behaviour is well represented by an approximation that uses physical insights into spin glasses.

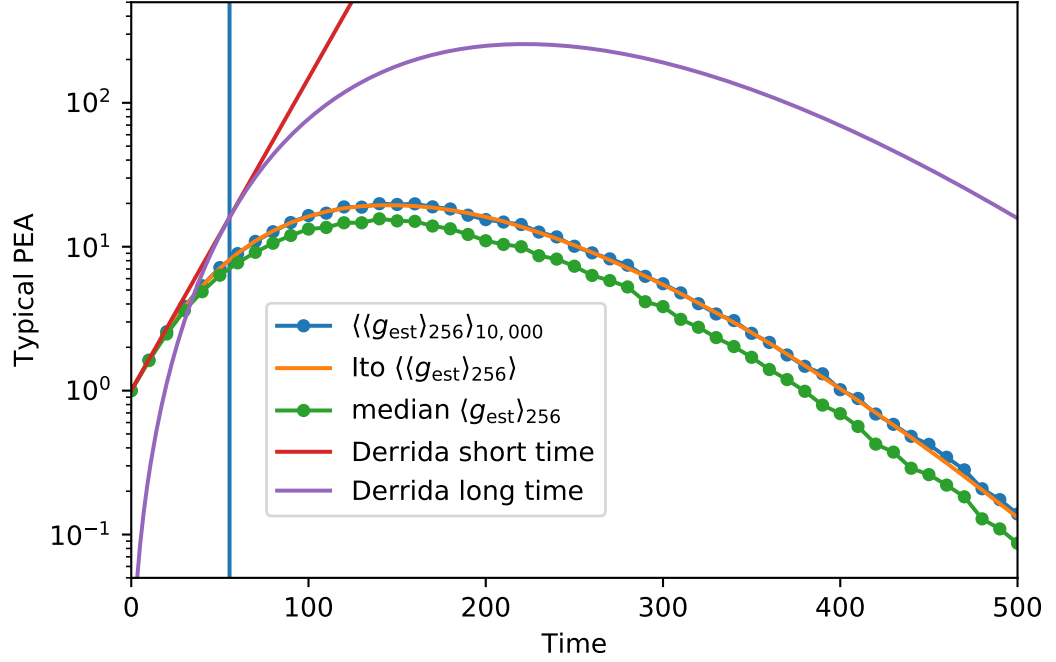


Figure 6.2: Lines are obtained by exponentiating the various exponential growth rates. **Blue line:** $\langle g_{\text{est}} \rangle_{256, 10,000}$ is the numerical mean (approximation of the expectation value) over a super-ensemble of $S = 10,000$ samples of g_{est} estimated in sub-ensembles of $N = 256$ GBMs each. **Green line:** median in a super-ensemble of S samples of g_{est} , each estimated in sub-ensembles of size N . **Yellow line:** An exact expression for $d \langle \ln \langle x \rangle_N \rangle$, derived using Itô calculus, see [47]. We evaluate the expression by Monte Carlo, and integrate, $\langle \ln \langle x \rangle_N \rangle = \int_0^t d \langle \ln \langle x \rangle_N \rangle$. Exponentiation yields the yellow line. **Red line:** short-time behavior, based on the random energy model, (Eq. 6.64). **Purple line:** long-time behavior, based on the random energy model, (Eq. 6.65). **Vertical line:** Crossover between the regimes at $t_c = \frac{2 \ln N}{\sigma^2}$, corresponding to $\beta_c = \frac{2(\ln 2)^{1/2}}{J}$. **Parameters:** $N = 256$, $S = 10,000$, $\mu = 0.05$, $\sigma = \sqrt{0.2}$.

{fig:1}

Chapter 7

Interactions

{chapter:Interactions

We add interactions to our null model of a population of noisy multipliers. It turns out that our decision criterion generates interesting emergent behaviour – cooperation, the sharing and pooling of resources, is often time-average growth optimal. This provides answers to the puzzles of why people cooperate and why we see socio-economic structure from the formation of firms to nation states with taxation and redistribution systems. We also ask whether assumptions about ergodicity in standard treatments of these topics are justified.

7.1 Cooperation

{section:Cooperation}

Under multiplicative growth, fluctuations are undesirable because they reduce time-average growth rates. In the long run, wealth $x_1(t)$ with noise term σ_1 will outperform wealth $x_2(t)$ with a larger noise term $\sigma_2 > \sigma_1$, in the sense that

$$\bar{g}(x_1) > \bar{g}(x_2) \quad (7.1)$$

with probability 1.

For this reason it is desirable to reduce fluctuations. One protocol that achieves this is pooling and sharing of resources. In Sec. 6.1 we explored the world created by the model of independent GBMs. This is a world where everyone experiences the same long-term growth rate. We want to explore the effect of the invention of cooperation. It turns out that cooperation increases growth rates, and this is a crucial insight.

Suppose two individuals, $x_1(t)$ and $x_2(t)$ decide to meet up every Monday, put all their wealth on a table, divide it in two equal amounts, and go back to their business (where their business is making their wealths follow GBM!) How would this operation affect the dynamic of the wealth of these two individuals?

Consider a discretised version of (Eq. 6.3), such as would be used in a numerical simulation. The non-cooperators grow according to

$$\Delta x_i(t) = x_i(t) \left[\mu \Delta t + \sigma \sqrt{\Delta t} \xi_i \right], \quad (7.2) \quad \{\text{eq:discrete_nonc_grow}\}$$

$$x_i(t + \Delta t) = x_i(t) + \Delta x_i(t), \quad (7.3) \quad \{\text{eq:discrete_nonc_coop}\}$$

where ξ_i are standard normal random variates, $\xi_i \sim \mathcal{N}(0, 1)$.

We imagine that the two previously non-cooperating entities, with resources $x_1(t)$ and $x_2(t)$, cooperate to produce two entities, whose resources we label $x_1^c(t)$ and $x_2^c(t)$ to distinguish them from the non-cooperating case. We envisage equal sharing of resources, $x_1^c = x_2^c$, and introduce a cooperation operator, $\oplus$, such that

$$x_1 \oplus x_2 = x_1^c + x_2^c. \quad (7.4)$$

In the discrete-time picture, each time step involves a two-phase process. First there is a growth phase, analogous to (Eq. 6.3), in which each cooperator increases its resources by

$$\Delta x_i^c(t) = x_i^c(t) \left[\mu \Delta t + \sigma \sqrt{\Delta t} \xi_i \right]. \quad (7.5) \quad \{\text{eq:discrete_coop_grow}\}$$

This is followed by a cooperation phase, replacing (Eq. 7.3), in which resources are pooled and shared equally among the cooperators:

$$x_i^c(t + \Delta t) = \frac{x_1^c(t) + \Delta x_1^c(t) + x_2^c(t) + \Delta x_2^c(t)}{2}. \quad (7.6) \quad \{\text{eq:discrete_coop_coop}\}$$

With this prescription both cooperators and their sum experience the following dynamic:

$$(x_1 \oplus x_2)(t + \Delta t) = (x_1 \oplus x_2)(t) \left[1 + \left(\mu \Delta t + \sigma \sqrt{\Delta t} \frac{\xi_1 + \xi_2}{2} \right) \right]. \quad (7.7) \quad \{\text{eq:discrete_cooperate}\}$$

For ease of notation we define

$$\xi_{1 \oplus 2} = \frac{\xi_1 + \xi_2}{\sqrt{2}}, \quad (7.8)$$

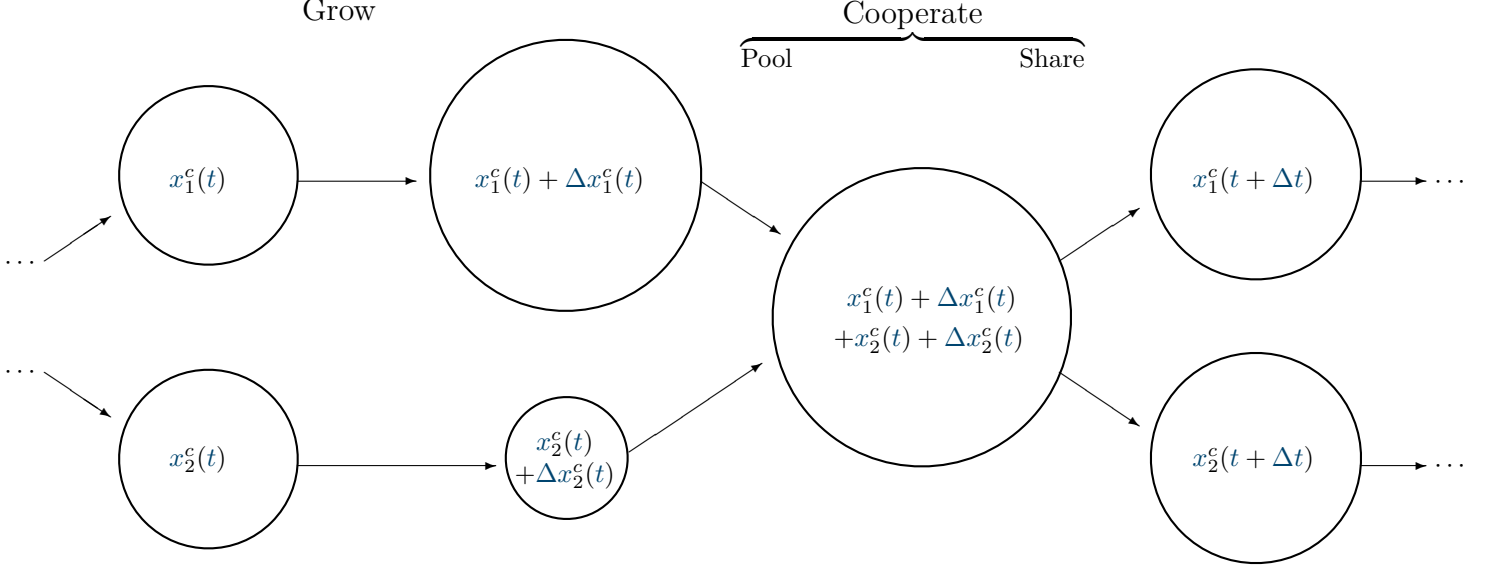


Figure 7.1: Cooperation dynamics. Cooperators start each time step with equal resources, then they *grow* independently according to (Eq. 7.5), then they *co-operate* by *pooling* resources and *sharing* them equally, then the next time step begins.

{fig:dynamics}

which is another standard Gaussian, $\xi_{1\oplus 2} \sim \mathcal{N}(0, 1)$. Letting the time increment $\Delta t \rightarrow 0$ we recover an equation of the same form as (Eq. 6.3) but with a different fluctuation amplitude,

$$d(x_1 \oplus x_2) = (x_1 \oplus x_2) \left(\mu dt + \frac{\sigma}{\sqrt{2}} dW_{1\oplus 2} \right). \quad (7.9)$$

The expectation values of a non-cooperator, $\langle x_1(t) \rangle$, and a corresponding cooperator, $\langle x_1^c(t) \rangle$, are identical. Based on expectation values, we thus cannot see any benefit of cooperation. Worse still, immediately after the growth phase, the better-off entity of a cooperating pair, $x_1^c(t_0) > x_2^c(t_0)$, say, would increase its expectation value from $\frac{x_1^c(t_0) + x_2^c(t_0)}{2} \exp(\mu(t - t_0))$ to $x_1^c(t_0) \exp(\mu(t - t_0))$ by breaking the cooperation. But it would be foolish to act on the basis of this analysis: the short-term gain from breaking cooperation is a one-off, and is dwarfed by the long-term multiplicative advantage of continued cooperation. An analysis based on expectation values finds that there is no reason for cooperation to arise, and that if it does arise there are good reasons for it to end, *i.e.* it will be fragile. Because expectation values are inappropriately used to evaluate future prospects, the observation of widespread cooperation constitutes a conundrum.

The solution of the conundrum comes from considering the time-average growth rate. The non-cooperating entities grow at $g_t(x_i) = \mu - \frac{\sigma^2}{2}$, whereas the cooperating unit benefits from a reduction of the amplitude of relative fluctuations and grows at $g_t(x_1 \oplus x_2) = \mu - \frac{\sigma^2}{4}$, and we have

$$g_t(x_1 \oplus x_2) > g_t(x_i) \quad (7.10)$$

for any non-zero noise amplitude. Imagine a world where cooperation does not exist, just like in Sec. 6.1. Now introduce into this world two individuals who have invented cooperation – very quickly this pair of individuals will become vastly more wealthy than anyone else. To keep up, others will have to start cooperating. The effect is illustrated in Fig. 7.2 by direct simulation of (Eq. 7.2)–(Eq. 7.3) and (Eq. 7.7).

Imagine again the pair of cooperators outperforming all of their peers. Other entities will have to form pairs to keep up, and the obvious next step is for larger cooperating units to form – groups of 3 may form, pairs of pairs, cooperation clusters of N individuals, and the larger the cooperating group the closer the time-average growth rate will get to the expectation value. For N cooperators, $x_1 \oplus x_2 \dots \oplus x_N$ the spurious drift term is $-\frac{\sigma^2}{2N}$, so that the time-average growth approaches expectation-value growth for large N . The approach to this upper bound as the number of cooperators increases favours the formation of social structure.

We may generalise to different drift terms, μ_i , and noise amplitudes, σ_i , for different individual entities. Whether cooperation is beneficial in the long run for any given entity depends on these parameters as follows. Entity 1 will benefit from cooperation with entity 2 if

$$\mu_1 - \frac{\sigma_1^2}{2} < \frac{\mu_1 + \mu_2}{2} - \frac{\sigma_1^2 + \sigma_2^2}{8}. \quad (7.11)$$

We emphasize that this inequality may be satisfied also if the expectation value of entity 1 grows faster than the expectation value of entity 2, *i.e.* if $\mu_1 > \mu_2$. An analysis of expectation values, again, is utterly misleading: the benefit conferred on entity 1 due to the fluctuation-reducing effect of cooperation may outweigh the cost of having to cooperate with an entity with smaller expectation value.

We may also generalise to correlations between the fluctuations experienced by different entities. These are uncorrelated in our model: the dW_i in (Eq. 6.3) and, consequently, the ξ_i in (Eq. 7.2) onwards are independent random variables. In reality, cooperators are often spatially localised and experience similar environmental conditions. By allowing correlations between the ξ_i , our model can be adapted to describe such situations. This is more technically difficult and we won't present it here. If you are interested, you can find the details in [46].

Notice the nature of the Monte-Carlo simulation in Fig. 7.2. No ensemble is constructed. Only individual trajectories are simulated and run for a time that is long enough for statistically significant features to rise above the noise. This method teases out of the dynamics what happens over time. The significance of any observed structure – its epistemological meaning – is immediately clear: this is what happens over time for an individual system (a cell, a person's wealth, *etc.*). Simulating an ensemble and averaging over members to remove noise does not tell the same story. The resulting features may not emerge over time. They are what happens on average in an ensemble, but – at least for GBM – this is not what happens to the individual with probability 1. For instance the pink dashed line in Fig. 7.2 is the ensemble average of $x_1(t)$, $x_2(t)$, and $(x_1 \oplus x_2)(t)/2$, and it has nothing to do with what happens in the individual trajectories over time.

When judged on expectation values, the apparent futility of cooperation is unsurprising because expectation values are the result for infinitely many

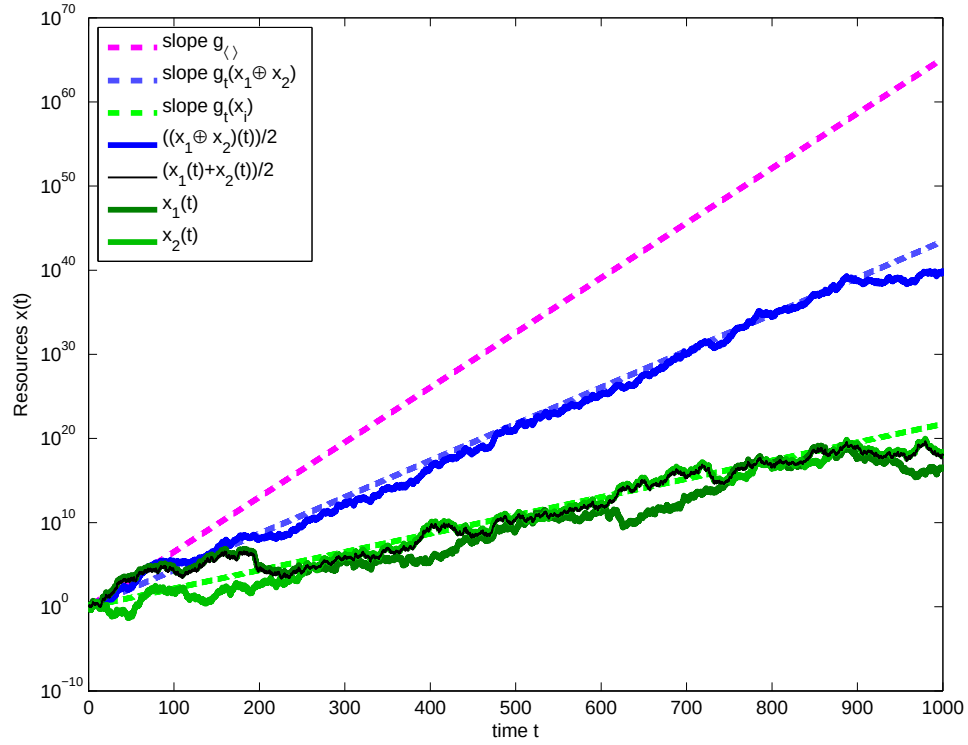


Figure 7.2: Typical trajectories for two non-cooperating (green) entities and for the corresponding cooperating unit (blue). Over time, the noise reduction for the cooperator leads to faster growth. Even without effects of specialisation or the emergence of new function, cooperation pays in the long run. The black thin line shows the average of the non-cooperating entities. While in the logarithmic vertical scale the average traces the more successful trajectory, it is far inferior to the cooperating unit. In a very literal mathematical sense the whole, $(x_1 \oplus x_2)(t)$, is more than the sum of its parts, $x_1(t) + x_2(t)$. The algebra of cooperation is not merely that of summation.

{fig:cooperate}

cooperators, and adding further cooperators cannot improve on this.

In our model the advantage of cooperation, and hence the emergence of social structure in the broadest sense – is purely a non-linear effect of fluctuations – cooperation reduces the magnitude of fluctuations, and over time (though not in expectation) this implies faster growth.

Another generalisation is partial cooperation – entities may share only a proportion of their resources, resembling taxation and redistribution. We discuss this in the next section.

7.2 Reallocation

{section:reallocation}

7.2.1 Introduction

{section:RGBM_intro}

In Sec. 6.1 we created a model world of independent trajectories of GBM. We studied how the distribution of the resulting random variables evolved over time. We saw that this is a world of broadening distributions, increasing inequality, and wealth condensation.

We introduced cooperation to it in Sec. 7.1 and saw how this increases the time-average growth rate for those who pool and share all of their resources. In this section we study what happens if a large number of individuals pool and share only a fraction of their resources. This is reminiscent of the taxation and redistribution – which we shall call “reallocation” – carried out by populations in the real world.

We will find that, while full cooperation between two individuals increases their growth rates, sufficiently fast reallocation from richer to poorer in a large population has two related effects. Firstly, everyone’s wealth grows in the long run at a rate close to that of the expectation value. Secondly, the distribution of rescaled wealth converges over time to a stable form. This means that, while wealth can still be distributed quite unequally, wealth condensation and the divergence of inequality no longer occur in our model. Of course, for this to be an interesting finding, we will have to quantify what we mean by “sufficiently fast reallocation.”

We will also find that when reallocation is too slow or, in particular, when it goes from poorer to richer – which we will label negative reallocation – no stable wealth distribution exists. In the latter case, the population splits into groups with positive and negative wealths, whose magnitudes grow exponentially.

Finally, having understood how our model behaves in each of these reallocation regimes, we will fit the model parameters to historical wealth data from the real world, specifically the United States. This will tell us which type of model behaviour best describes the dynamics of the US wealth distribution in both the recent and more distant past. You might find the results surprising – we certainly did!

7.2.2 The ergodic hypothesis in economics

{section:RGBM_EH}

Of course, we are not the first to study resource distributions and inequality in economics. This topic has a long history, going back at least as far as Vilfredo Pareto’s work in the late 19th century [40] (in which he introduced the power-law distribution we discussed in Sec. 6.1.7). It’s worth spending a few moments

considering how the subject is approached classically, before we say how we approach it.

Economists studying such distributions usually make the following assumption in their models: that the distributions converge in the long run to a unique and stable form, regardless of initial conditions. This allows them to study the stable distribution, for which many statistical techniques exist, and to ignore the transient phenomena preceding it, which are far harder to analyse. Paul Samuelson called this the “ergodic hypothesis” [53, pp. 11-12]. It’s easy to see why: if this convergence happens, then the time average of the observable under study will equal its ensemble average with respect to the stable distribution.¹

Economics is often concerned with growth and a growing quantity cannot be ergodic in Samuelson’s sense, because its distribution never stabilises. This suggests the simplifying ergodic hypothesis should *never* be made. Not so fast! Although rarely stated, a common strategy to salvage these techniques is to find a transformation of the non-ergodic process that produces a meaningful ergodic observable. If such an ergodic observable can be derived, then classical analytical techniques may still be used. We have already seen in the context of gambles that expected utility theory can be viewed as transformation of non-ergodic wealth increments into ergodic utility increments. Expectation values, which would otherwise be misleading, then quantify time-average growth of the decision-maker’s wealth.

Studies of wealth distributions also employ this strategy. Individual wealth is modelled as a growing quantity. Dividing by the population average transforms this to a rescaled wealth, as in Sec. 6.1.5, which is hypothesised to be ergodic. For example, [4, p. 130] “impose assumptions ... that guarantee the existence and uniqueness of a limit stationary distribution.” The idea is to take advantage of the simplicity with which the stable distribution can be analysed, *e.g.* to predict the effects of policies encoded in model parameters.

There is, however, an elephant in the room. To our knowledge, the validity of the ergodic hypothesis for rescaled wealth has never been tested empirically. It’s certainly invalid for the GBM model world we studied previously because, as we saw in Sec. 6.1.5, rescaled wealth has an ever-broadening log-normal distribution. That doesn’t seem to say much, as most reasonable people would consider our model world – containing a population of individuals whose wealths multiply noisily and who never interact – a tad unrealistic. The model we are about to present will not only extend our understanding from this simple model world to one containing interactions, but also will allow us to test the hypothesis. This is because it has regimes, *i.e.* combinations of parameters, for which rescaled wealth is and isn’t ergodic. This contrasts with models typically used by economists, which have the ergodic hypothesis “baked in.”

If it is reasonable to assume a stable distribution exists, we must also consider how long convergence to would take after a change of parameters. It’s no use if convergence in the model takes the whole of human history, if we are using it to estimate the effect of a tax policy over the next election cycle. Therefore, treating a stable model distribution as representative of an empirical wealth distribution implies an assumption of fast convergence. As the late Tony Atkinson pointed out, “the speed of convergence makes a great deal of difference to the

¹Convergence to a unique and stable distribution is a sufficient but not necessary condition for an ergodic observable, as we have defined it.

way in which we think about the model” [?]. We will use our model to discuss this point. Without further ado, let us introduce it.

7.2.3 Reallocating geometric Brownian motion

{section:RGBM_model}

Our model, called [reallocating geometric Brownian motion \(RGBM\)](#), is a system of N individuals whose wealths, $x_i(t)$, evolve according to the stochastic differential equation,

$$dx_i = x_i [(\mu - \varphi)dt + \sigma dW_i(t)] + \varphi \langle x \rangle_N dt, \quad (7.12) \quad \{\text{eq:rgbm}\}$$

for $i = 1 \dots N$. In effect, we have added to [GBM](#) a simple reallocation mechanism. Over a time step, dt , each individual pays a fixed proportion of its wealth, $\varphi x_i dt$, into a central pot (“contributes to society”) and gets back an equal share of the pot, $\varphi \langle x \rangle_N dt$, (“benefits from society”). We can think of this as applying a wealth tax, say of 1% per year, to everyone’s wealth and then redistributing the tax revenues equally. Note that the reallocation parameter, φ , is, like μ , a rate with dimensions per unit time. Note also that when $\varphi = 0$, we recover our old friend, [GBM](#), in which individuals grow their wealths without interacting.

[RGBM](#) is our null model of an exponentially growing economy with social structure. It is intended to capture only the most general features of the dynamics of wealth. A more complex model would treat the economy as a system of agents that interact with each other through a network of relationships. These relationships include trade in goods and services, employment, taxation, welfare payments, using public infrastructure (roads, schools, a legal system, social security, scientific research, and so on), insurance, wealth transfers through inheritance and gifts, and so on. It would be a hopeless task to list exhaustively all these interactions, let alone model them explicitly. Instead we introduce a single parameter – the reallocation rate, φ – to represent their net effect. If φ is positive, the direction of net reallocation is from richer to poorer. If negative, it is from poorer to richer.

We will see shortly that [RGBM](#) has both ergodic and non-ergodic regimes, characterised to a good approximation by the sign of φ . $\varphi > 0$ produces an ergodic regime, in which wealths are positive, distributed with a Pareto tail, and confined around their mean value. $\varphi < 0$ produces a non-ergodic regime, in which the population splits into two classes, characterised by positive and negative wealths which diverge away from the mean.

We offer a couple of health warnings. In [RGBM](#), like in [GBM](#), there are no additive changes akin to labour income and consumption. This is unproblematic for large wealths, where additive changes are dwarfed by capital gains. For small wealths, however, wages and consumption are significant and empirical distributions look rather different for low and high wealths [17]. We modelled earnings explicitly in [?] and found this didn’t generate insights different from [RGBM](#) when we fit both models to real wealth data. We note also, as [?, p. 41] put it, that our agents “do not marry or have children or die or even grow old.” Therefore, the individual in our setup is best imagined as a household or a family, *i.e.* some long-lasting unit into which personal events are subsumed.

Having specified the model, we will use insights from Sec. 6.2 to understand how rescaled wealth is distributed in the ergodic and non-ergodic regimes. Then we will show briefly our results from fitting the model to historical wealth data

from the United States. The full technical details of this fitting exercise are beyond the scope of these notes – if you are interested, you can find “chapter and verse” in [?]. Fitting φ to data will allow us to answer the important questions:

- What is the net reallocating effect of socio-economic structure on the wealth distribution?
- Are observations consistent with the ergodic hypothesis that the rescaled wealth distribution converges to a stable distribution?
- If so, how long does it take, after a change in conditions, for the rescaled wealth distribution to reach the stable distribution?

7.2.4 Model behaviour

{section:RGBM_behaviour}

It is instructive to write (Eq. 7.12) as

$$dx_i = \underbrace{x_i [\mu dt + \sigma dW_i(t)]}_{\text{Growth}} - \underbrace{\varphi(x_i - \langle x \rangle_N) dt}_{\text{Reallocation}}. \quad (7.13) \quad \{\text{eq:rgbm_ou}\}$$

This resembles GBM with a mean-reverting term like that of [60] in physics and [62] in finance. It exposes the importance of the sign of φ . We discuss the two regimes in turn.

Positive φ

For $\varphi > 0$, individual wealth, $x_i(t)$, reverts to the sample mean, $\langle x(t) \rangle_N$. We explored some of the properties of sample mean in Sec. 6.2 for wealths undergoing GBM. In particular, we saw that a short-time (or large-sample or low-volatility) self-averaging regime exists, (Eq. 6.59) $t < t_c \equiv \frac{2 \ln N}{\sigma^2}$, where the sample mean is approximated well by the ensemble average,

$$\langle x(t) \rangle_N \sim \langle x(t) \rangle = \exp(\mu t). \quad (7.14) \quad \{\text{eq:rgbm_self}\}$$

(The final equality assumes, as previously, that $x_i(0) = 1$ for all i .) It turns out that the same self-averaging approximation can be made for wealths undergoing RGBM, (Eq. 7.12), when the reallocation rate, φ , is above some critical threshold:

$$\varphi > \varphi_c \equiv \frac{\sigma^2}{2 \ln N}. \quad (7.15) \quad \{\text{eq:tau_c}\}$$

Showing this is technically difficult [8] and we will confine ourselves to sketching the key ideas in Sec. 7.2.5 below. It won’t have escaped your attention that $\varphi_c = t_c^{-1}$ and, indeed, you will shortly have an intuition for why.

Fitting the model to data yields parameter values for which φ_c is extremely small. For example, typical parameters for US wealth data are $N = 10^8$ and $\sigma = 0.2 \text{ year}^{-1/2}$, giving $\varphi_c = 0.1\% \text{ year}^{-1}$ (or $t_c = 900$ years). Accounting for the uncertainty in the fitted parameters makes this statistically indistinguishable from $\varphi_c = 0$.

This means we can safely make the self-averaging approximation for the entire positive φ regime. That’s great news, because it means we can rescale wealth by the ensemble average, $\langle x(t) \rangle = \exp(\mu t)$, as we did in Sec. 6.1.5 for GBM, and

not have to worry about pesky finite N effects. Following the same procedure as there gives us a simple SDE in the rescaled wealth, $y_i(t) = x_i(t) \exp(-\mu t)$:

$$dy_i = y_i \sigma dW_i(t) - \varphi(y_i - 1) dt. \quad (7.16) \quad \{\text{eq:rgbm_ou_re}\}$$

Note that the common growth rate, μ , has been scaled out as it was in Sec. 6.1.5.

The distribution of $y_i(t)$ can be found by solving the corresponding Fokker-Planck equation, which we will do in Sec. 7.2.5. For now, we will just quote the result: a stable distribution exists with a power-law tail, to which the distribution of rescaled wealth converges over time. The distribution has a name – the Inverse Gamma Distribution – and a probability density function:

$$\mathcal{P}(y) = \frac{(\zeta - 1)^\zeta}{\Gamma(\zeta)} e^{-\frac{\zeta-1}{y}} y^{-(1+\zeta)}. \quad (7.17) \quad \{\text{eq:disti}\}$$

$\zeta = 1 + 2\varphi/\sigma^2$ is the Pareto tail index (corresponding to $\alpha - 1$ in Sec. 6.1.7) and $\Gamma(\cdot)$ is the gamma function.

Example forms of the stationary distribution are shown in Figure 7.3. The usual stylised facts are recovered: the larger σ (more randomness in the returns) and the smaller φ (less social cohesion), the smaller the tail index ζ and the fatter the tail of the distribution. Fitted φ values give typical ζ values between 1 and 2 for the different datasets analysed, consistent with observed tail indices between 1.2 to 1.6 (see [?] for details). Not only does RGBM predict a realistic functional form for the distribution of rescaled wealth, but also it admits fitted parameter values which match observed tails. The inability to do this is a known weakness of earnings-based models (again, see [?] for discussion).

For positive reallocation, (Eq. 7.16) and extensions of it have received much attention in statistical mechanics and econophysics [10, 9]. As a combination of GBM and a mean-reverting process it is a simple and analytically tractable stochastic process. [?] provide an overview of the literature and known results.

Negative φ

For $\varphi < 0$ the model exhibits mean repulsion rather than reversion. The ergodic hypothesis is invalid and no stationary wealth distribution exists. The population splits into those above the mean and those below the mean. Whereas in RGBM with non-negative φ it is impossible for wealth to become negative, negative φ leads to negative wealth. No longer is total economic wealth a limit to the wealth of the richest individual because the poorest develop large negative wealth. The wealth of the rich in the population increases exponentially away from the mean, and the wealth of the poor becomes negative and exponentially large in magnitude, see Figure 7.4.

Such splitting of the population is a common feature of non-ergodic processes. If rescaled wealth were an ergodic process, then individuals would, over long enough time, experience all parts of its distribution. People would spend 99 percent of their time as “the 99 percent” and 1 percent of their time as “the 1 percent”. Therefore, the social mobility implicit in models that assume ergodicity might not exist in reality if that assumption is invalid. That inequality and immobility have been linked [13, 29, 5] may be unsurprising if both are viewed as consequences of non-ergodic wealth or income.

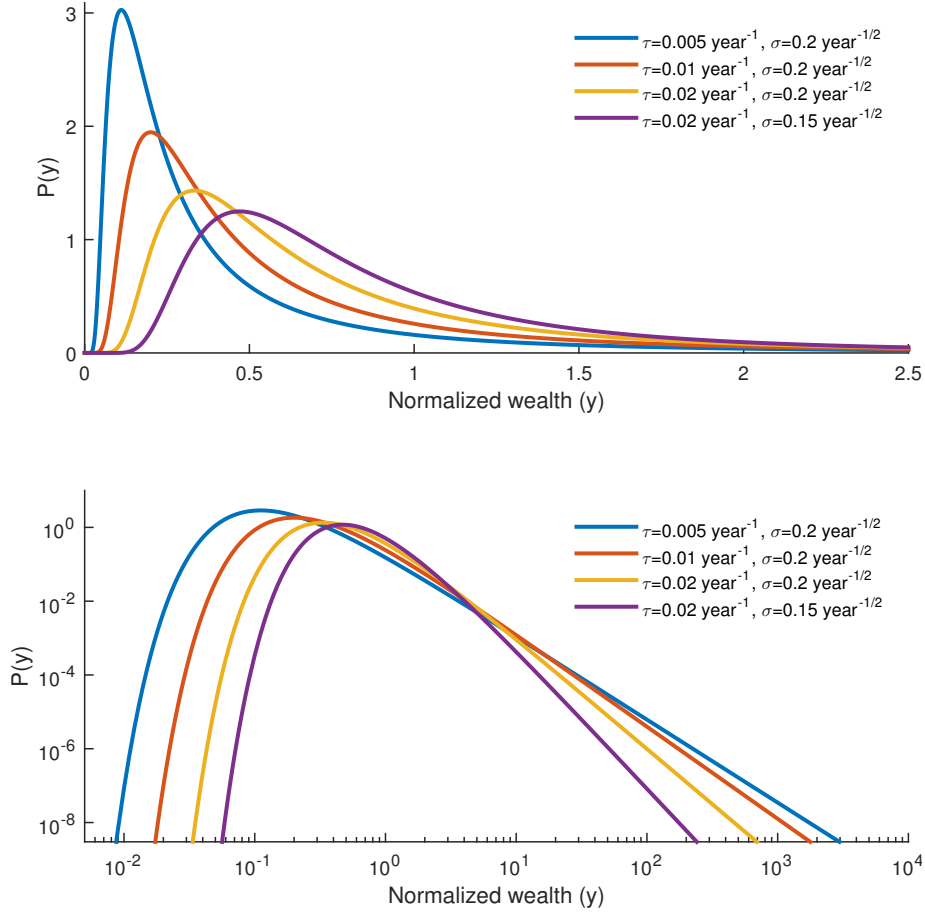


Figure 7.3: The stationary distribution for RGBM with positive φ . Top – linear scales; Bottom – logarithmic scales.

{fig:dist}

7.2.5 Derivation of the stable distribution

{section:RGBM_stable}

In this section we will sketch the argument for why we can make the self-averaging approximation, (Eq. 7.14), in RGBM with sufficiently fast positive reallocation, (Eq. 7.15). This is shown rigorously in [8]. Then we will solve the Fokker-Planck equation for the rescaled wealth and derive the inverse gamma distribution, (Eq. 7.17).

We presented arguments in Sec. 6.2.1 for why wealth in GBM is self-averaging, $\langle x(t) \rangle_N \sim \langle x(t) \rangle = \exp(\mu t)$ for short time. By mapping from GBM to the random energy model in Sec. 6.2.2, we showed that “short time” means $t < t_c$, where $t_c = 2 \ln N / \sigma^2$. We can think of this as follows: t_c is the timescale over which the inequality-increasing effects of noisy multiplicative growth drive wealths apart, such that a finite sample of wealths stops self-averaging and becomes dominated by a few trajectories.

Let’s now think about what happens when we add reallocation to GBM, creating RGBM. φ is the reallocation rate, so φ^{-1} is reallocation timescale, *i.e.* the timescale over which the inequality-reducing effects of reallocation pull wealths

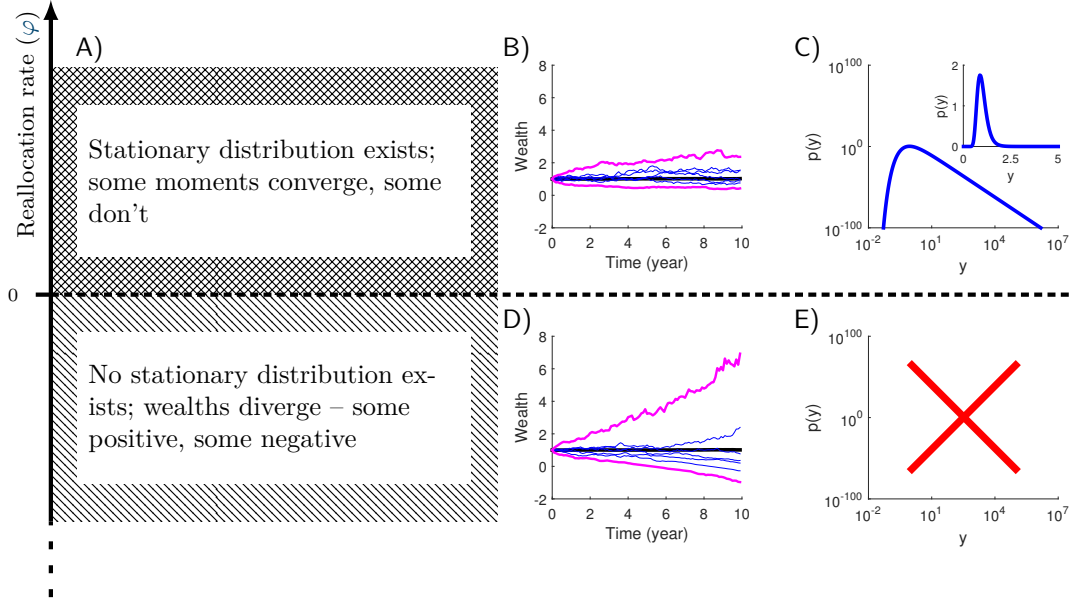


Figure 7.4: Regimes of RGBM. A) $\varphi = 0$ separates the two regimes of RGBM. For $\varphi > 0$, a stationary wealth distribution exists. For $\varphi < 0$, no stationary wealth distribution exists and wealths diverge. B) Simulations of RGBM with $N = 1000$, $\mu = 0.021 \text{ year}^{-1}$ (presented after rescaling by $\exp(\mu t)$), $\sigma = 0.14 \text{ year}^{-1/2}$, $x_i(0) = 1$, $\varphi = 0.15 \text{ year}^{-1}$. Magenta lines: largest and smallest wealths, blue lines: five randomly chosen wealth trajectories, black line: sample mean. C) The stationary distribution to which the system in B) converges. Inset: same distribution on linear scales. D) Similar to B), with $\varphi = -0.15 \text{ year}^{-1}$. E) in the $\varphi < 0$ regime, no stationary wealth distribution exists.

{fig:regimes}

together. If $\varphi^{-1} > t_c$, then reallocation happens too slowly to prevent the expiry of self-averaging. However, if $\varphi^{-1} < t_c$, then reallocation pulls wealths together more quickly than they get driven apart, continually “resetting” the sample and allowing self-averaging to be maintained indefinitely. Converting this condition into a reallocation rate, we get $\varphi > t_c^{-1}$, as in (Eq. 7.15). As mentioned in Sec. 7.2.4, this becomes indistinguishable from $\varphi > 0$ for realistic parameters, so the self-averaging approximation can be made safely for all positive φ when using the model to study real economies.

We can now approximate the rescaled wealth, $y_i(t) = x_i(t)/\langle x(t) \rangle_N$, as $y_i(t) = x_i(t)\exp(-\mu t)$. Using Itô’s formula as we did Sec. 6.1.5 yields the following SDE for rescaled wealth, now under RGBM instead of GBM:

$$dy = \sigma y dW - \varphi (y - 1) dt. \quad (7.18)$$

This is an Itô equation with drift term $A = \varphi(y - 1)$ and diffusion term $B = y\sigma$. Such equations imply ordinary second-order differential equations that describe the evolution of the PDF, called Fokker-Planck equations. The Fokker-Planck equation describes the change in probability density, at any point in (rescaled wealth) space, due to the action of the drift term (like advection in a fluid) and

due to the diffusion term (like heat spreading). In this case, we have

$$\frac{d\mathcal{P}(y,t)}{dt} = \frac{\partial}{\partial y} [A\mathcal{P}(y,t)] + \frac{1}{2} \frac{\partial^2}{\partial y^2} [B^2\mathcal{P}(y,t)]. \quad (7.19)$$

The steady-state Fokker-Planck equation for the PDF, $\mathcal{P}(y)$, is obtained by setting the time derivative to zero,

$$\frac{\sigma^2}{2} (y^2\mathcal{P})_{yy} + \varphi[(y-1)\mathcal{P}]_y = 0. \quad (7.20) \quad \{\text{eq:fokker_planck}\}$$

Positive wealth subjected to continuous-time multiplicative dynamics with non-negative reallocation can never reach zero. Therefore, we solve Equation (7.20) with boundary condition $\mathcal{P}(0) = 0$ to give

$$\mathcal{P}(y) = C(\zeta) e^{-\frac{\zeta-1}{y}} y^{-(1+\zeta)}, \quad (7.21)$$

where

$$\zeta = 1 + \frac{2\varphi}{\sigma^2} \quad (7.22)$$

and

$$C(\zeta) = \frac{(\zeta-1)^\zeta}{\Gamma(\zeta)}, \quad (7.23)$$

with the gamma function $\Gamma(\zeta) = \int_0^\infty x^{\zeta-1} e^{-x} dx$. The distribution has a power-law tail as $y \rightarrow \infty$, resembling Pareto's oft-confirmed observation that the frequency of large wealths tends to decay as a power-law [40]. The exponent of the power-law, ζ , is called the Pareto parameter and is one measure of economic inequality.

7.2.6 Moments and convergence times

{section:RGBM_moments}

The inverse gamma distribution, (Eq. 7.17), has a power-law tail. This means that, for positive reallocation, while some of the lower moments of the stable rescaled wealth distribution may exist, higher moments will not. Specifically, the k^{th} moment diverges if $k > \zeta$.

If we find parameters consistent with positive reallocation when we fit our model to data, we will be interested in whether certain statistics – such as the variance – exist. We will also want to know how long it takes the distribution to converge sufficiently to its stable form for them to be meaningful. Here we derive a condition for the convergence of the variance and calculate its convergence time, noting also the general procedure for other statistics.

The variance of y is a combination of the first moment, $\langle y \rangle$ (the ensemble average), and the second moment, $\langle y^2 \rangle$:

$$V(y) = \langle y^2 \rangle - \langle y \rangle^2 \quad (7.24)$$

Thus we need to find $\langle y \rangle$ and $\langle y^2 \rangle$ in order to determine the variance.

The first moment of the rescaled wealth is, by definition, $\langle y \rangle = 1$. To find the dynamic of the second moment, we start with the SDE for the rescaled wealth,

$$dy = \sigma y dW - \varphi(y-1) dt, \quad (7.25) \quad \{\text{eq:rescaledSDE}\}$$

and follow a now familiar procedure. We insert $f(y, t) = y^2$ into Itô's formula,

$$df = \frac{\partial f}{\partial t} dt + \frac{\partial f}{\partial y} dy + \frac{1}{2} \frac{\partial^2 f}{\partial y^2} dy^2 \quad (7.26)$$

to obtain

$$d(y^2) = 2y dy + dy^2. \quad (7.27) \quad \{\text{eq:diff2}\}$$

We substitute (Eq. 7.27) for dy to get terms at orders dW , dt , dW^2 , dt^2 , and $dWdt$. The scaling of BM allows us to replace dW^2 by dt and we ignore terms at $o(dt)$. This yields

$$d(y^2) = 2\sigma y^2 dW - (2\varphi - \sigma^2) y^2 dt + 2\varphi y dt. \quad (7.28)$$

Taking expectations on both sides and noting that $\langle y \rangle = 1$ gives us an ordinary differential equation for the second moment:

$$\frac{d\langle y^2 \rangle}{dt} = -(2\varphi - \sigma^2) \langle y^2 \rangle + 2\varphi \quad (7.29) \quad \{\text{eq:avediff2}\}$$

with solution

$$\langle y(t)^2 \rangle = \frac{2\varphi}{2\varphi - \sigma^2} + \left(\langle y(0)^2 \rangle - \frac{2\varphi}{2\varphi - \sigma^2} \right) e^{-(2\varphi - \sigma^2)t}. \quad (7.30) \quad \{\text{eq:avediff3}\}$$

The variance $V(t) = \langle y(t)^2 \rangle - 1$ therefore follows

$$V(t) = V_\infty + (V_0 - V_\infty) e^{-(2\varphi - \sigma^2)t}, \quad (7.31) \quad \{\text{eq:var1}\}$$

where V_0 is the initial variance and

$$V_\infty = \frac{2\varphi}{2\varphi - \sigma^2}. \quad (7.32) \quad \{\text{eq:varinf}\}$$

V converges in time to the asymptote, V_∞ , provided the exponential in (Eq. 7.31) is decaying. This can be expressed as a condition on φ

$$\varphi > \frac{\sigma^2}{2}, \quad (7.33)$$

which is the same as the condition we noted previously for the second moment to exist: $\zeta > k$ where $k = 2$.

Clearly, for negative values of φ the condition cannot be satisfied, and the variance (and inequality) of the wealth distribution will diverge. In the regime where the variance exists, $\varphi > \sigma^2/2$, it also follows from (Eq. 7.31) that the convergence time of the variance is $1/(2\varphi - \sigma^2)$.

As φ increases, increasingly high moments of the distribution converge to finite values. The above procedure for finding the second moment (and thereby the variance) can be applied to the k^{th} moment, just by changing the second power y^2 to y^k . Therefore, any other cumulant can be found as a combination of the relevant moments. For instance, [?] also compute the third cumulant.

7.2.7 Fitting United States wealth data

{section:RGBM_data}

We have introduced the RGBM model and understood its basic properties. It is a simple model of an interacting population of noisy multiplicative growers. We expect it to be more realistic than GBM ($\varphi = 0$) because we know that in the real world people interact. In particular, large populations have over centuries developed public institutions and infrastructure, to which everyone contributes and from which everyone benefits. At first glance, therefore, we might expect to find that RGBM with positive φ to fit real wealth data better than GBM or, indeed, RGBM with negative φ .

Additionally, if this is true and if associated convergence times are shorter than the timescales of policy changes, it would indicate that the ergodic hypothesis is warranted and a helpful modelling assumption. If not, then the hypothesis would be unjustified and could be acting as a serious constraint on models of economic inequality, generating misleading analyses and recommendations.

In [?] we fit the RGBM model to historical wealth data from the United States for the last hundred years. We won't include full technical description of this empirical analysis here. It would be too long and our main aim is to communicate the ideas we use to think about problems and build models. If you want to know more, please read the paper (and let us know what you think!) However, the results are interesting and, to us at least, a little shocking, so we include a brief summary.

The basic setup is to fix the values of μ , σ , and N in our RGBM model using data about, respectively, aggregate economic growth, stock market volatility, and population data; and then to find, by numerical simulation of (Eq. 7.12), the time series of $\varphi(t)$ values which best reproduces historical wealth shares in the United States. The wealth share, $S_q(t)$, is a type of inequality measure. It is defined as the proportion of total wealth owned by the richest fraction q of the population. So, for example, $S_{0.1} = 0.8$ means that the richest 10% of the population own 80% of the total wealth. Reproducing the historical wealth shares is one way of reproducing approximately the level of inequality in the wealth distribution, and the nice thing is that economists Emmanuel Saez and Gabriel Zucman have estimated around a century's worth of wealth shares for the United States [?].

Fitting the model to these data will address two main questions:

- Is the ergodic hypothesis valid for rescaled wealth in the United States? For it to be valid, fitted values of $\varphi(t)$ must be robustly positive.
- If $\varphi(t)$ is robustly positive, is convergence of the distribution to its stable form fast enough for the distribution to be used as a representative of the empirical wealth distribution?

Note that we have relaxed the model slightly. φ is a fixed parameter in (Eq. 7.12) but we allow it to vary with time in our empirical analysis.

Figure 7.5 (top) shows the results of fitting the RGBM model to the wealth share data in [?]. There are large annual fluctuations in $\varphi_q(t)$ (black line) but we are more interested in longer-term changes in reallocation driven by structural economic and political changes. To show these, we smooth the data by taking a central 10-year moving average, $\tilde{\varphi}_q(t)$ (red line), where the window is truncated at the ends of the time series. We also show the uncertainty in this moving average (red envelope).

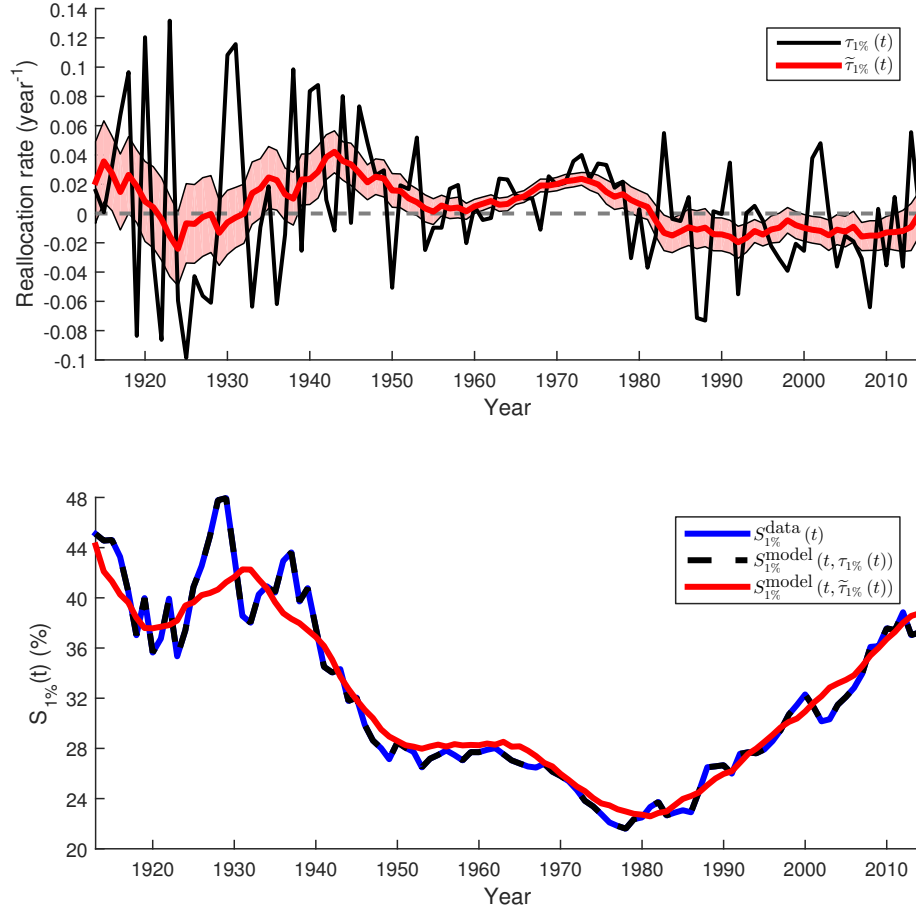


Figure 7.5: Fitted effective reallocation rates. Calculations done using $\mu = 0.021 \text{ year}^{-1}$ and $\sigma = 0.16 \text{ year}^{-1/2}$. Top: $\varphi_{1\%}(t)$ (black) and $\tilde{\varphi}_{1\%}(t)$ (red). Translucent envelopes indicate one standard error in the moving averages. Bottom: $S_{1\%}^{\text{data}}(t)$ (blue), $S_{1\%}^{\text{model}}(t, \tau_{1\%}(t))$ (dashed black), based on the 10-year moving average $\tilde{\varphi}_{1\%}(t)$ (red).

{fig:tau}

To ensure the smoothing does not introduce artificial biases, we reverse the procedure and use $\tilde{\varphi}_q(t)$ to propagate the initially inverse gamma-distributed wealths and determine the wealth shares $S_q^{\text{model}}(t)$. The good agreement with $S_q^{\text{data}}(t)$ suggests that the smoothed $\tilde{\varphi}_q(t)$ is meaningful, see Figure 7.5 (bottom).

The “take home” is this: while the effective reallocation rate, $\tilde{\varphi}(t)$, was positive for most of the last hundred years, it was negative – *i.e.* reallocation from poorer to richer – from the mid-1980s onwards. Furthermore, even when $\varphi(t)$ was positive, associated convergence times (estimated both by numerical simulation and by plugging fitted model parameters into the results of Sec. 7.2.6) were very long compared to the typical times between policy changes – from several decades to several centuries. This makes the answer to both our questions above a resounding “No.”

What shocked us most when we first encountered these results was the existence of long periods with negative reallocation. We began the analysis imagining that GBM (*i.e.* $\varphi = 0$, no interactions) was a really crazy model of the real world. After all, we personally pay our taxes and use public services, and the London Mathematical Laboratory is supported by generous charitable reallocations! We imagined that we would see $\varphi > 0$ but that it might be so small that convergence times would be too long for the ergodic hypothesis to be useful. Instead we found that, recently in the United States at least, reallocation has been consistently negative. In our model, this corresponds to wealths being driven apart, populations splitting into groups with positive and negative wealths, and no convergence to a stable distribution of rescaled wealth.

In retrospect, perhaps we shouldn't have been so surprised. Qualitatively, our results echo the findings that the rich are experiencing higher growth rates of their wealth than the poor [50, ?] and that the cumulative wealth of the poorest 50 percent of the American population was negative during 2008–2013 [?, 57].

The economic phenomena that trouble theorists most – such as diverging inequality, social immobility, and the emergence of negative wealth – are difficult to reproduce in models that make the ergodic hypothesis. In our simple model, this is easy to see: in the ergodic regime, $\varphi > 0$, our model cannot reproduce these phenomena at all. One may be tempted to conclude that their existence is a sign of special conditions prevailing in the real world – collusion and conspiracies. But if we admit the possibility of non-ergodicity, $\varphi \leq 0$, it becomes clear that these phenomena can easily emerge in an economy that does not actively guard against them.

Chapter 8

Markets

{chapter:Markets}

This lecture applies the ideas developed in the preceding lectures to markets.

We set up a simple portfolio selection problem in a market of two assets: one risky, like shares; and the other riskless, like a bank deposit. We ask how an investor would best allocate his money between the two assets, which we phrase in terms of his leverage. We review the classical approach, which can't answer this question without additional information about the investor's risk preferences. We then use the decision theory we've developed so far to answer the question unambiguously, by deriving the optimal leverage which maximises the investment's time-average growth rate.

This is a Gedankenexperiment that will lead us to predict and the empirically discover certain regularities in the price statistics of freely traded assets. We consider what this objectively defined optimal leverage might mean for financial markets themselves. If all the participants in a market aim for the same optimal leverage, does this constrain the prices and price fluctuations that emerge from their trading? We argue that it does, and we quantify how. Our treatment of the problem generates a theory of noise in stock prices, resolves the so-called equity premium puzzle (and the price volatility puzzle), provides a natural framework for setting central-bank interest rates, and even suggests a method for fraud detection. We test our predictions using data collected from the American and German stock markets as well as Bitcoin and Bernie Madoff's Ponzi scheme.

8.1 Optimal leverage

8.1.1 A model market

We consider assets whose values follow multiplicative dynamics, which we will model using **GBM**. In general, an amount x invested in such an asset evolves according to the **SDE**,

$$dx = x(\mu dt + \sigma dW), \quad (8.1) \quad \{\text{eq:sde_gbm}\}$$

where μ is the drift and σ is the volatility. By now we are very familiar with this equation and how to solve it.

To keep things simple, we imagine a market of two assets. One asset is riskless: the growth in its value is known deterministically and comes with a cast-iron guarantee.¹ This might correspond in reality to a bank deposit. The other asset is risky: there is uncertainty over what its value will be in the future. This might correspond to a share in a company or a collection of shares in different companies. We will think of it simply as stock.

An amount x_0 invested in the riskless asset evolves according to

$$dx_0 = x_0 \mu_r dt. \quad (8.2) \quad \{\text{eq:sde_0}\}$$

μ_r is the riskless drift, known in finance as the riskless rate of return.² There is no volatility term. In effect, we have set $\sigma = 0$. We know with certainty what x_0 will be at any point in the future:

$$x_0(t_0 + \Delta t) = x_0(t_0) \exp(\mu_r \Delta t). \quad (8.3) \quad \{\text{eq:sde_0_soln}\}$$

An amount x_1 invested in the risky asset evolves according to

$$dx_1 = x_1(\mu_s dt + \sigma_s dW), \quad (8.4) \quad \{\text{eq:sde_1}\}$$

where $\mu_s > \mu_r$ is the risky drift and $\sigma_s > 0$ is the volatility (the subscript s stands for stock). μ_s is also known in finance as the expected return.³ This equation has solution

$$x_1(t_0 + \Delta t) = x_1(t_0) \exp \left[\left(\mu_s - \frac{\sigma_s^2}{2} \right) \Delta t + \sigma_s W(\Delta t) \right], \quad (8.5) \quad \{\text{eq:sde_1_soln}\}$$

which is a random variable.

We will refer to the difference

$$\mu_e = \mu_s - \mu_r \quad (8.6) \quad \{\text{eq:def_mue}\}$$

as the excess drift.

¹Such guarantees are easy to offer in a model. In the real world, one should be very suspicious of anything that comes with a “cast-iron guarantee”.

²In general we will eschew financial terminology for rates. Economics has failed to define them clearly, with the result that different quantities, like $\bar{g}$ and $g_{\square}$, are often conflated. The definitions developed in these lectures are aimed at avoiding such confusion.

³Probably because it's the growth rate of the expected value, see (Eq. 8.13).

8.1.2 Leverage

Let's turn to the concept of leverage. Imagine a very simple portfolio of value x_ℓ , out of which ℓx_ℓ is invested in stock and the remainder, $(1 - \ell)x_\ell$, is put in the bank. ℓ is known as the leverage. It is the fraction of the total investment assigned to the risky asset. $\ell = 0$ corresponds to a portfolio consisting only of bank deposits. $\ell = 1$ corresponds to a portfolio only of stock.

You would be forgiven for thinking that prudence dictates $0 \leq \ell \leq 1$, *i.e.* that we invest some of our money in stock and keep the rest in the bank. However, the financial markets have found all sorts of exciting ways for us to invest almost any amount in an asset. For example, we can make $\ell > 1$ by borrowing money from the bank to buy more stock than we could have bought with only our own money.⁴ We can even make $\ell < 0$ by borrowing stock (a negative investment in the risky asset), selling it, and putting the money raised in the bank. In the financial world this practice is called short selling.

Each investment in our portfolio experiences the same relative fluctuations as the asset in which it has been made. The overall change in the portfolio's value is, therefore,

$$dx_\ell = (1 - \ell)x_\ell \frac{dx_0}{x_0} + \ell x_\ell \frac{dx_1}{x_1}. \quad (8.7)$$

Substituting in (Eq. 8.2) and (Eq. 8.4) gives the SDE for a leveraged investment in the risky asset,

$$dx_\ell = x_\ell[(\mu_r + \ell\mu_e)dt + \ell\sigma_s dW], \quad (8.8) \quad \{\text{eq:sde_1}\}$$

with solution,

$$x_\ell(t_0 + \Delta t) = x_\ell(t_0) \exp \left[\left(\mu_r + \ell\mu_e - \frac{\ell^2 \sigma_s^2}{2} \right) \Delta t + \ell\sigma_s W(\Delta t) \right]. \quad (8.9) \quad \{\text{eq:sde_1_soln}\}$$

We can now see why we labelled investments in the riskless and risky assets by x_0 and x_1 : when $\ell = 0$, x_ℓ follows the same evolution as x_0 ; and when $\ell = 1$, it evolves as x_1 .

In our model ℓ , once chosen, is held constant over time. This means that our model portfolio must be continuously rebalanced to ensure that the ratio of stock to total investment stays fixed at ℓ . For example, imagine our stock investment fluctuates up a lot over a short time-step, while our bank deposit only accrues a little interest. Immediately we have slightly more than ℓ of the portfolio's value in stock, and slightly less than $1 - \ell$ of its value in the bank. To return the leverage to ℓ , we need to sell some stock and deposit the proceeds in the bank, Fig. 8.1. In (Eq. 8.8) we are imagining that this happens continuously.⁵

8.1.3 Portfolio theory

{section:pf_theory}

Our simple model portfolio parametrised by ℓ allows us to ask the

⁴This doesn't immediately affect the portfolio's value. The bank loan constitutes a negative investment in the riskless asset, whose value cancels the value of the stock we bought with it. Of course, the change in the portfolio's composition will affect its future value.

⁵In reality, of course, that's not possible. We could try to get close by rebalancing frequently. However, every time we buy or sell an asset in a real market, we pay transaction costs, such as broker's fees and transaction taxes. This means that frequent rebalancing in the real world can be costly.

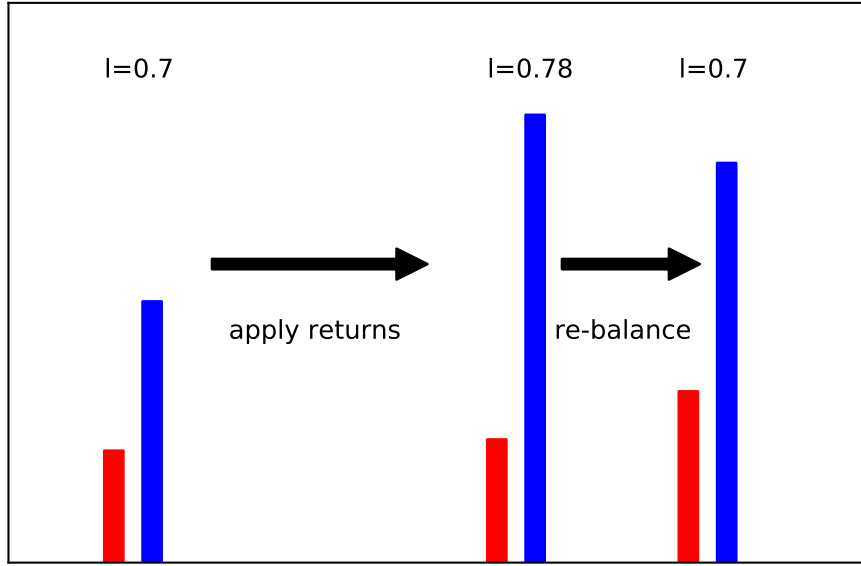


Figure 8.1: Investment in a risky (blue) and riskless (red) asset. Leverage starts at $\ell = 0.7$, meaning 70% is invested in the risky asset, and 30% in the riskless. The investment then experiences the relative returns of the market. In this case the risky asset goes up quite a bit, and the riskless asset goes up less. Consequently, the leverage changes. To maintain a constant leverage, some risky asset has to be sold in return for some riskless, a step known as re-balancing. fig:rebalance

Question:

What is the optimal value of ℓ ?

This is similar to choosing between gambles, for which we have already developed a decision theory. The main difference is that we are now choosing from a continuum of gambles, each characterised by a value of ℓ , whereas previously we were choosing between discrete gambles. The principle, however, is the same: we will maximise the time-average growth rate of our investment.

Before we do this, let's review the classical treatment of the problem so that we appreciate the wider context. Intuitively, people understand there is some kind of trade-off between risk and reward. In our model of a generic multiplicative asset, (Eq. 8.1), we could use σ as a proxy for risk and μ as a proxy for reward. Ideally we want an investment with large μ and small σ , but we also acknowledge the rule-of-thumb that assets with larger μ tend to have larger σ .⁶ This is why we model our risky asset as having a positive excess return, $\mu_e > 0$, over the riskless asset.

Intuition will only take us so far. A rigorous treatment of the portfolio selection problem was first attempted by Markowitz in 1952 [31]. He suggested defining a portfolio with parameters (σ_i, μ_i) as efficient if there exists no rival portfolio with parameters (σ_j, μ_j) with $\mu_j \geq \mu_i$ and $\sigma_j \leq \sigma_i$. Markowitz argued that it is unwise to invest in a portfolio which is not efficient.

⁶A “no such thing as a free lunch” type of rule.

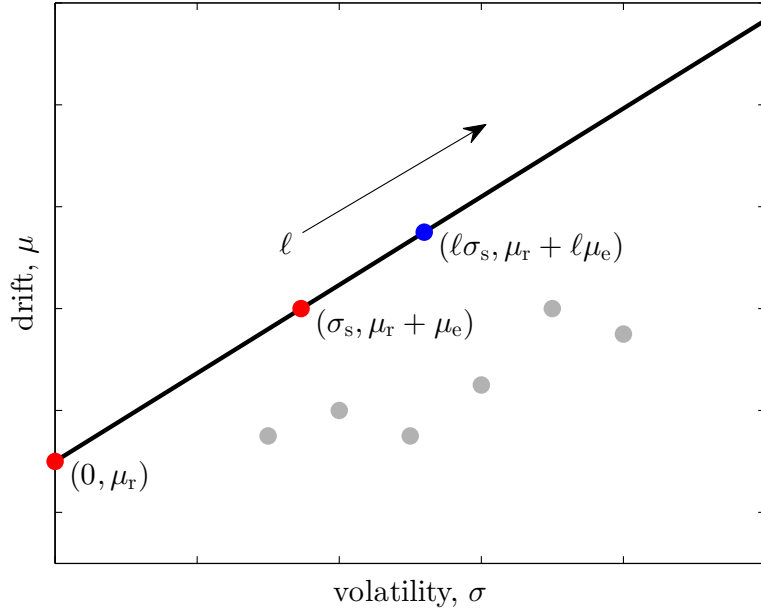


Figure 8.2: The Markowitz portfolio selection picture. The red dots are the locations in the (σ, μ) -plane of portfolios containing only the riskless (left) or risky (right) asset. The blue dot is one possible leveraged portfolio, in this case with $\ell > 1$. All possible leveraged portfolios lie on the black line, (Eq. 8.10). The grey dots are hypothetical alternative portfolios, containing different assets excluded from our simple portfolio problem. Their location below and to the right of the black line makes them inefficient under the Markowitz scheme. fig:markowitz

In our problem, we are comparing portfolios with parameters $(\ell\sigma_s, \mu_r + \ell\mu_e)$. These lie on a straight line in the (σ, μ) -plane,⁷

$$\mu = \mu_r + \left(\frac{\mu_e}{\sigma_s} \right) \sigma, \quad (8.10) \quad \{\text{eq:frontier}\}$$

shown schematically in Fig. 8.2. Under Markowitz’s classification, all of the leveraged portfolios on this line are efficient.⁸ Therefore, any leverage we choose gives a portfolio the classification would recommend. This does not help answer our question. By itself, Markowitz’s approach is agnostic to leverage: it requires additional information to select a specific portfolio as the optimum. Markowitz was aware of this limitation and argued that the optimal portfolio could be identified by considering the investor’s risk preferences [33]: “The proper choice among portfolios depends on the willingness and ability of the investor to assume risk.” The ultimate reliance on personal preferences is a theme that runs through all of economic theory, and we have seen examples of it in previous chapters. Of course, economic theory, or finance theory, is not meant to replace human judgement and imagination, nor is it supposed to force certain products and behaviours on people – but it is supposed to guide decisions. A statement like “this is the action that will eventually make your wealth grow fastest, assuming we have chosen models and estimated parameters correctly” is helpful and leaves plenty of room for judgement calls. As in decision theory in general,

⁷Derived by eliminating ℓ from the equations $\mu = \mu_r + \ell\mu_e$ and $\sigma = \ell\sigma_s$.

⁸Indeed, in finance this line is called the efficient frontier.

risk preferences are typically included in the portfolio selection problem through a utility function, and thereby the theory of finance inherits a major cornerstone of economics, along with all its problems.

8.1.4 Sharpe ratio

{section:sharpe}

The Sharpe ratio [56] for an asset with drift μ and volatility σ is defined as

$$S \equiv \frac{\mu - \mu_r}{\sigma} \quad (8.11) \quad \{\text{eq:def_sharpe}\}$$

It is the gradient of the straight line in the (σ, μ) -plane which passes through the riskless asset and the asset in question. S is often used as a convenient shorthand for applying Markowitz's ideas, since choosing the portfolio with the highest S from the set of available portfolios is equivalent to choosing an efficient portfolio. The optimal leverage problem, however, important as it is, cannot be resolved by this procedure. The reason is this: all of our leveraged portfolios lie on the same line, (Eq. 8.10), and so all of them have the same Sharpe ratio, which is simply the line's gradient:

$$S_\ell = \frac{\mu_e}{\sigma_s}. \quad (8.12) \quad \{\text{eq:sharpe_1}\}$$

This is insensitive to the leverage ℓ , resulting in the same non-advice as the Markowitz approach. Sharpe also suggested considering risk preferences to resolve the optimal portfolio: "The investor's task is to select from among the efficient portfolios the one that he considers most desirable, based on his particular feelings regarding risk and expected return" [56]. Once more we encounter the theme of reliance on personal preferences in finance and economics. It is no surprise that this switch from formal recommendation (a portfolio should be on the efficient frontier/have a high Sharpe ratio) to personal preference (pick the one whose risks you feel most comfortable with) happens too soon, from our perspective. In general, we use a formal model as long as it's useful, and when we have pushed it to the point where it yields no further reliable answers we switch to personal judgement. The conceptual superiority of our analysis (relying on averages over time not over parallel worlds) allows us to push the model a little further, and the point where we previously had to switch now seems premature.

8.1.5 Expected return

{section:exp_ret}

We noted previously that the growth rate of the expectation value of the risky price is the risky drift, μ_s , also known as the expected return. This is because

$$\langle x_1(t_0 + \Delta t) \rangle = \langle x_1(t_0) \rangle \exp(\mu_s \Delta t), \quad (8.13) \quad \{\text{eq:exp_ret}\}$$

which, as a multiplicative process, has growth rate

$$g_m(\langle x_1 \rangle) = \frac{\Delta \ln \langle x_1 \rangle}{\Delta t} = \mu_s. \quad (8.14)$$

It follows immediately from comparison of (Eq. 8.4) and (Eq. 8.8) that the expected value of the leveraged portfolio grows at

$$g_m(\langle x_\ell \rangle) = \mu_r + \ell \mu_e. \quad (8.15) \quad \{\text{eq:exp_ret_1}\}$$

This illustrates why a portfolio theory which is insensitive to leverage, such as that of Markowitz and Sharpe⁹, is potentially dangerous. Because it doesn't alert the investor to the dangers of over-leveraging and considers any leverage optimal, the investor must find an additional criterion. There is nothing in the formalism that would prevent the investor to select as that additional criterion the expected return in (Eq. 8.15). Consequently he would maximise his leverage, $\ell \rightarrow \infty$. This, as we will shortly see, would almost surely ruin him.

8.1.6 Growth rate maximisation

We now understand the classical approach, as applied to a very simple portfolio problem, and we are aware of its limitations. What does our own decision theory have to say?

The time-average growth rate of the leveraged portfolio is

$$\bar{g}_m(\ell) \equiv \lim_{\Delta t \rightarrow \infty} \{g_m(x_\ell, \Delta t)\} = \lim_{\Delta t \rightarrow \infty} \left\{ \frac{\Delta \ln x_\ell}{\Delta t} \right\}. \quad (8.16) \quad \{\text{eq:g_1_def}\}$$

Note that (Eq. 8.1) fully defines the dynamics – finance works with stochastic processes that include the mode of repetition. In contrast, classical decision theory works with gambles where the appropriate mode of repetition has to be guessed. This is a key advantage of finance over economics, and it means that while from the utility perspective the problem of portfolio selection is under-specified (we need to guess the utility function), from the perspective of time optimization it is not. We have all the information we need, the problem is well posed. The time-average growth rate will depend on ℓ . Inserting the expression for x_ℓ from (Eq. 8.9) in (Eq. 8.16) gives

$$\bar{g}_m(\ell) = \lim_{\Delta t \rightarrow \infty} \left\{ \frac{1}{\Delta t} \left[\left(\mu_r + \ell \mu_e - \frac{\ell^2 \sigma_s^2}{2} \right) \Delta t + \ell \sigma_s W(\Delta t) \right] \right\}, \quad (8.17) \quad \{\text{eq:g_1_noisy}\}$$

which, since $W(\Delta t)/\Delta t \sim \Delta t^{-1/2} \rightarrow 0$ as $\Delta t \rightarrow \infty$, converges to

$$\bar{g}_m(\ell) = \mu_r + \ell \mu_e - \frac{\ell^2 \sigma_s^2}{2}. \quad (8.18) \quad \{\text{eq:g_1_quadratic}\}$$

This is a quadratic in ℓ with an unambiguous maximum¹⁰ at

$$\ell_{\text{opt}} = \frac{\mu_e}{\sigma_s^2}. \quad (8.19) \quad \{\text{eq:l_opt}\}$$

ℓ_{opt} is the optimal leverage which defines the portfolio with the highest time-average growth rate. Classical theory is indifferent to where on the line in Fig. 8.2 we choose to be. Our decision theory, however, selects a particular point on that line,¹¹ which answers the question we posed at the start of Sec. 8.1.3. This is the key result. We note that all we need to know are the parameters μ_r , μ_s , and σ_s of the two assets in our market. In particular, it is defined

⁹Both recipients of the 1990 Alfred Nobel Memorial Prize in Economic Sciences.

¹⁰Derived, for example, by setting $\frac{d\bar{g}_m}{d\ell} = 0$.

¹¹Subsequently Markowitz became aware of this point, which he called the “Kelly-Latané” point in [32], referring to [24, 26].

his leverage in either direction will, if he is able to apply enough leverage, lose wealth. Indeed, if his leveraging ability is unlimited, he will find to his horror what is easily seen in (Eq. 8.18), that $\bar{g}_m(\ell)$ diverges negatively as $\ell \rightarrow \pm\infty$. He will lose fast.

8.2 Stochastic market efficiency

{section:Stochastic}

8.2.1 A fundamental measure

{section:fundamental}

Aside from being insensitive to leverage, the Sharpe ratio, $S = \mu_e/\sigma_s$, is a dimensionful quantity. Its unit is $(\text{time unit})^{-1/2}$. This means that its numerical value is arbitrary (since it depends on the choice of time unit) and tells us nothing fundamental about the system under study. For example, a portfolio with $S = 5$ per square root of one year has $S = 5(365)^{-1/2} \approx 0.26$ per square root of one day. Same portfolio, different numbers.

The optimal leverage, $\ell_{\text{opt}} = \mu_e/\sigma_s^2$, which differs from the Sharpe ratio by a factor of $1/\sigma_s$, is a dimensionless quantity. Therefore, its numerical value does not depend on choice of units and has the potential to carry fundamental information about the system.¹² We could view ℓ_{opt} as the fundamental measure of a portfolio's quality, similarly to how S is viewed in the classical picture. A portfolio with a high optimal leverage must represent a good investment opportunity to justify such a large commitment of the investor's funds.

However, the significance of ℓ_{opt} runs deeper than this. The portfolio to which ℓ_{opt} refers is that which optimally allocates money between the risky and riskless assets in our model market. Therefore, it tells us much about conditions in that market and, by extension, in the wider model economy. A high ℓ_{opt} indicates an economic environment in which investors are incentivised to take risks. Conversely, a low or negative ℓ_{opt} indicates little incentive to take risks. This raises a tantalising

Question:

Are there any special numerical values of ℓ_{opt} which describe different market regimes, or to which markets are attracted?

8.2.2 Relaxing the model

Our model market contains assets whose prices follow GBM with constant drift and volatility parameters.¹³ Once specified, μ_r , μ_s , and σ_s are static and, therefore, so is ℓ_{opt} . This limits the extent to which we can explore the question, since we cannot consider changes in ℓ_{opt} . To make progress we need to relax the model. We must consider which parts of the model are relevant to the question, and which parts can be discarded without grave loss.

The GBM-based model is useful because it motivates the idea of an objectively optimal leverage which maximises the growth of an investment over time. It also provides an expression for ℓ_{opt} in terms of parameters which, in essence, describe the market conditions under which prices fluctuate. These parameters

¹²See Barenblatt's modern classic on scaling [3].

¹³A tautology, since (Eq. 8.1) only describes a GBM if μ and σ are constant in time.

have correspondences with quantities we can measure in real markets. All of this is useful.

However, we are ultimately interested in real markets, and their asset prices do not strictly follow GBM (because nothing in nature¹⁴ truly does). In particular, real market conditions are not static. They change over time, albeit on a longer time scale than that of the price fluctuations. In this context, the model assumption of constant μ_r , μ_s , and σ_s is restrictive and unhelpful. We will relax it and imagine a less constrained model market where these parameters, and therefore ℓ_{opt} , are allowed to vary slowly. We will not build a detailed mathematical formulation of this, but instead use the idea to run some simple thought experiments.

8.2.3 Efficiency

One way of approaching a question like this is to invoke the concept of efficiency. In economics this has a specific meaning in the context of financial markets, which we will mention momentarily. In general terms, an efficient system is one which is already well optimised and whose performance cannot be improved upon by simple actions. For example, a refrigerator is efficient if it maintains a cool internal environment while consuming little electrical power and emitting little noise¹⁵. Similarly, Markowitz’s portfolios were efficient because the investor could do no better than to choose one of them.

The “efficient market hypothesis” of classical economics treats markets as efficient processors of information. It claims that the price of an asset in an efficient market reflects all of the publicly available information about it. The corollary is that no market participant, without access to privileged information, can consistently beat the market simply by choosing the prices at which he buys and sells assets. We shall refer to this hypothesis as ordinary efficiency.¹⁶

We will consider a different sort of efficiency, where we think not about the price at which assets are bought and sold in our model market, but instead about the leverage that is applied to them. Let’s run a thought experiment.

Thought experiment: efficiency under leverage

Imagine that $\ell_{\text{opt}} > 1$ in our model market. This would mean that the simple strategy of borrowing money to buy stock will achieve faster long-run growth than buying stock only with our own money. If we associate putting all our money in stock, $\ell = 1$, with an investment in the market, then it would be a trivial matter for us to beat the market (by doing nothing more sophisticated than investing borrowed money).

¹⁴We consider markets to be part of nature and appropriate objects of scientific study.

¹⁵This is a subtle reference to a fridge whose presence once graced the offices of LML. It was not efficient.

¹⁶This is not the most precise hypothesis ever hypothesised. What does it mean for a price to “reflect” information? Presumably this involves some comparison between the observed price of the asset and its true value contingent on that information. But only the former is observable, while the latter evades clear definition. Similarly, what does it mean to “beat the market”? Presumably something to do with achieving a higher growth rate than a general, naïve investment in the overall market. But what investment, exactly? We will leave these legitimate questions unanswered here, since our focus is a different form of market efficiency. The interested reader can consult the comprehensive review in [55].

Similarly, imagine that $\ell_{\text{opt}} < 1$. In this scenario, the market could again be beaten very easily by leaving some money in the bank (and, if $\ell_{\text{opt}} < 0$, by short selling).

It would strain language to consider our market efficient if consistent out-performance were so straightforward to achieve. This suggests a different, fluctuations-based notion of market efficiency, which we call stochastic market efficiency: it is impossible for a market participant without privileged information to beat a stochastically efficient market simply by choosing the *amount* he invests in stock, *i.e.* by choosing his leverage.¹⁷ We believe real markets to be stochastically efficient. Therefore, we make the following

Hypothesis: stochastic market efficiency

Real markets self-organise such that

$$\ell_{\text{opt}} = 1$$

(8.21) {eq:sme_strong}

is an attractive point for their stochastic properties.

These stochastic properties are represented by μ_r , μ_s , and σ_s in the relaxed model that permits dynamic adjustment of their values.

8.2.4 Stability

Another approach to the question in Sec. 8.2.1, whose style owes more to physics than economics, is to consider the stability of the system under study (here, the market) and how this depends on the value of the observable in question (here, ℓ_{opt}).

Systems which persist over long time scales tend to be stable. Many of the systems we find in nature therefore include stabilising elements. Unstable systems tend to last for shorter times,¹⁸ so we observe them less frequently. With this in mind, let's think about what different values of ℓ_{opt} imply for systemic stability.

Thought experiment: stability under leverage

Imagine that $\ell_{\text{opt}} > 1$ in our relaxed model. Since it is an objectively optimal leverage which does not depend on investor idiosyncrasies, this means that *everyone* in the market should want to borrow money to buy stock. But, if that's true, who's going to lend the money and who's going to sell the stock?

Similarly, imagine that $\ell_{\text{opt}} < 0$. This means that *everyone* should want to borrow stock and sell it for cash. But, if that's true, who's going to lend

¹⁷This resembles ordinary efficiency except that we have replaced price by amount.

¹⁸We use the comparative "shorter" here. This does not mean short. It is quite possible for an unstable system to remain in flux for a long time in human terms, perhaps giving the illusion of stability. Indeed, much of classical economic theory is predicated on the idea that economies are at or close to equilibrium, *i.e.* stable. We would argue that economies are fundamentally far-from-equilibrium systems and must be modelled as such, even if their dynamics unfold over time scales much longer than our day-to-day affairs.

the stock and who's going to relinquish their cash to buy it?

Unless there are enough market participants disinterested in their time-average growth rate to take the unfavourable sides of these deals – which in our model we will assume there aren't – then neither of these situations is globally stable. It's hard to imagine them persisting for long before trading activity causes one or more of the market parameters to change, returning ℓ_{opt} to a stable value.

This thought experiment suggests that we could observe markets with $0 \leq \ell_{\text{opt}} \leq 1$. We've defined things so that the optimal leverage is the proportion of one's wealth ideally invested in the risk asset, *i.e.* in shares. $\ell_{\text{opt}} < 1$ probably doesn't correspond to a macroeconomically desirable state, in the sense that people would be incentivised to keep their money in bank accounts rather than to invest it. But it wouldn't be unstable purely from the leverage point of view. This hints at an interesting generalisation: what if both assets are risky? What if we're comparing two share indexes or two shares, eliminating cash from our considerations? The same stability arguments would suggest that the optimal leverage in such a portfolio – now defined as the proportion ideally held in one of the assets – should be in the range $0 \leq \ell_{\text{opt}} \leq 1$, which reflects the symmetry between two risky assets. Our original hypothesis then describes the special case where one asset does not fluctuate.

Hypothesis: stochastic market efficiency for general asset pairs

On sufficiently long time scales, the range $0 \leq \ell_{\text{opt}} \leq 1$ is an attractor for the stochastic properties of pairs of risky assets.

Whether in our original setup real markets spend significant time in a phase with $\ell_{\text{opt}} < 1$ is difficult to tell. Realistic optimal leverage depends significantly on such things as trading costs, access to credit and risky assets in lending programs, infrastructure and technology. All of these factors differ among market participants. They suggest that measuring optimal leverage in the simplest way – using market prices only and assuming zero trading costs – will produce an over-estimate, compared to the value experienced by real market participants who incur realistic costs.

The dynamical adjustment, or self-organisation, of the market parameters takes place through the trading activity of market participants. In particular, this creates feedback loops, which cause prices and fluctuation amplitudes to change, returning ℓ_{opt} to a stable value whenever it strays. To be truly convincing, we should propose plausible trading mechanisms through which these feedback loops arise. We do this in [?]. Since they involve details about how trading takes place in financial markets (in which we assume the typical attendee of these lectures is disinterested) we shall not rehearse them here. The primary drivers of our hypothesis are the efficiency and stability arguments we've just made.

Furthermore, there are additional reasons why we would favour the strong form of the hypothesis over long time scales. The main one is that an economy in which ℓ_{opt} is close to, or even less than, zero gives people no incentive to invest in productive business activity. Such an economy would appear paralysed, resembling perhaps those periods in history to which economists refer as

depressions. We'd like to think that economies are not systematically attracted to such states. The other reasons are more technical, to do with the different interest rates accrued on deposits and loans, and the costs associated with buying and selling assets. These are described in [?].

8.2.5 Prediction accuracy

{section:Prediction_accu

When we simply use observed price changes to compute changes in portfolio value, we will overestimate optimal leverage. Denoting an estimate from real data of ℓ_{opt} by $\ell_{\text{opt}}^{\text{est}}$, we will have $\ell_{\text{opt}}^{\text{est}} > \ell_{\text{opt}}$. This is because real investments incur transaction costs and other costs associated with leveraging. We therefore expect naïvely measured values to be a little too big, $\ell_{\text{opt}}^{\text{est}} > 1$.

Ignoring these effects for the moment, we can predict how close we expect observed optimal leverage to be to its predicted value, 1, when we estimate it from a finite time series, let's say of daily returns. Even assuming ideal conditions – that ℓ_{opt} really is attracted to 1 and that any effects that lead to systematic mis-estimation can be neglected, we expect random deviations from $\ell_{\text{opt}}^{\text{est}} = 1$ to increase as the time series gets shorter. To take an extreme example, with daily data, the observed optimal leverage over a window of one day does not exist. Either $\ell_{\text{opt}}^{\text{est}} \rightarrow +\infty$ if the risky return beats the deposit rate on that day; or $\ell_{\text{opt}}^{\text{est}} \rightarrow -\infty$ if deposits beat the risky asset. Indeed, the magnitude of the observed optimal leverage will diverge for any window over which the daily risky returns are either all greater than, or all less than, the daily returns on federal funds. This is unlikely for windows of months or years but using daily data it happens commonly over windows of days or weeks. Even without this divergence, shorter windows are more likely to result in larger positive and negative optimal leverages because relative fluctuations are larger over shorter time scales.

To quantify this idea we compute $g_{\text{m}}(x_{\ell}, \Delta t)$, *i.e.* the time-average growth rate observed in an investment following (Eq. 8.8) after a finite time Δt . That's simply the expression on the RHS of (Eq. 8.17) without taking the $\Delta t \rightarrow \infty$ limit:

$$g_{\text{m}}(x_{\ell}, \Delta t) = \mu_{\text{r}} + \ell \mu_{\text{e}} - \frac{(\ell \sigma_{\text{s}})^2}{2} + \frac{\ell \sigma_{\text{s}} W(\Delta t)}{\Delta t} \quad (8.22) \quad \{\text{eq:g_1_finite}\}$$

Maximizing this expression generates a noisy estimate for the optimal leverage over a window of length Δt :

$$\ell_{\text{opt}}^{\text{est}}(\Delta t, N = 1) = \ell_{\text{opt}} + \frac{W(\Delta t)}{\sigma_{\text{s}} \Delta t}. \quad (8.23)$$

Thus, in the model, optimal leverage for finite time series is normally distributed with mean ℓ_{opt} and standard deviation

$$\text{stdev}(\ell_{\text{opt}}^{\text{est}}(\Delta t)) = \frac{1}{\sigma_{\text{s}} \Delta t^{1/2}}. \quad (8.24) \quad \{\text{eq:variance}\}$$

We will use this quantity as the standard error for the prediction $\ell_{\text{opt}}^{\text{est}} \approx 1$.

8.3 Applications of stochastic market efficiency

{section:applications}

Our solution of the portfolio selection problem, and the concept of stochastic market efficiency that follows from it are momentous developments in the theory

of financial markets. While in a sense the puzzle pieces necessary for these developments have been around for a while, it seems that they haven't so far been put together in quite the way we've presented them. This is evident from many puzzles and questions that exist in finance and economics but which have a straight-forward solution in the conceptual space we have developed in these notes. Before we look at data in Sec. 8.4 to understand how accurately this fundamental theoretical work reflects empirical reality, we apply it to some of these famous puzzles to illustrate the type of solution the theory generates.

8.3.1 A theory of noise – Sisyphus discovers prices

{section:a_theory}

According to stochastic market efficiency, prices of risky assets must fluctuate if an excess drift exists, $\mu_e > 0$, simply because the market would otherwise become unstable. But, if price fluctuations are necessary for stability, then the intellectual basis for price efficiency – that changes in price are driven by the arrival of new economic information – cannot be the whole truth. Or –depending on what is meant by “information” – it may be an empty circular statement. At least some component of observed fluctuations must be driven by the leverage feedbacks described in section Sec. 8.2, which enforce leverage efficiency and which have little to do with the arrival of meaningful economic information.

Black differentiated between information-based and other types of price fluctuation, referring to the latter as “noise” and regarding it as a symptom of inaccurate information and market *in*-efficiency [7]. However, substituting $\ell_{\text{opt}} = 1$ in (Eq. 8.19) yields

$$(\sigma_s)^2 = \mu_e. \quad (8.25) \quad \{\text{eq:noise}\}$$

Purely based on stochastic market efficiency, prices must fluctuate, and we can even quantify by how much. An asset whose expectation value grows faster than the value of the riskless asset must fluctuate; otherwise systemic stability will be undermined. This is a radical departure from conventional thinking, and it has practical consequences. For instance, prices “discovered” at ever higher trading frequencies must necessarily reveal more ups and downs, but this noise is self-generated, imposed by the requirement of leverage stability. That stability is the genesis of volatility constitutes a theory of noise requiring no appeal to the arrival of unspecified information, accurate or not.

8.3.2 Solution of the equity premium puzzle

{section:solution}

The term “equity premium” has been used to describe the “premium” I receive for holding “equity” – meaning the extra growth rate my wealth experiences, compared to riskless interest rates, if I invest it in shares, or some other risky asset. We denote this with the symbol π^m (for premium).

Puzzles usually arise because we look at things from the wrong angle, with a mental model that doesn't reflect reality. The equity premium puzzle is no exception. The researchers who first studied the equity premium had a model in mind whereby equity prices are set based on the consumption preferences of the population [35]. We won't go into these models in detail because they're the wrong angle. But the story is this: consumption-based asset pricing models come to the conclusion that no one should ever hold cash. In these models only people who are pathologically terrified of losing money would hold cash,

and such people just don't exist. But if you assume that they don't exist, then the models would predict a very different equity premium. So something is not working. In 2016 LeRoy summarized the state of the debate as follows [28]: “Most analysts believe that no single convincing explanation has been provided for the volatility of equity prices. The conclusion that appears to follow from the equity premium and price volatility puzzles is that, for whatever reason, prices of financial assets do not behave as the theory of consumption-based asset pricing predicts.”

Of course the difference between the time-average growth rates of the risky asset ($\ell = 1$) and the riskless asset ($\ell = 0$) can be measured. It is

$$\pi^m \equiv \bar{g}_m(1) - \bar{g}_m(0) \quad (8.26) \quad \{\text{eq:epdef}_1\}$$

$$= \mu_e - \frac{(\sigma_s)^2}{2}. \quad (8.27) \quad \{\text{eq:epdef}\}$$

What value would we expect the equity premium to take in a real market? Substituting (Eq. 8.25) into (Eq. 8.27), it follows that the equity premium is attracted to

$$\pi^m = \frac{(\sigma_s)^2}{2}. \quad (8.28) \quad \{\text{eq:epval}\}$$

It is important to remember that any theoretical prediction comes with a band of uncertainty around it. The next step is, therefore, to estimate the uncertainty in the equity premium. We do this by considering finite measurement times and substitute (Eq. 8.22) into (Eq. 8.26)

$$\text{stdev}(\pi^m(\Delta t)) = \frac{\sigma_s}{\Delta t^{1/2}}. \quad (8.29) \quad \{\text{eq:pim_error}\}$$

In Sec. 8.4 we will use this as our definition of one statistical standard error. The tension may by now have become unbearable, so here's a sneak preview of what's to come empirically: Our estimate of $(\sigma_s)^2$ for the [Deutscher Aktienindex \(German stock index\) \(DAX\)](#) since its inception in 1987 is 5.3%*p.a.*. Half of that is 2.6%*p.a.*, and the standard error is 4.2% *p.a.*, meaning the 2-standard-error range around the predicted value of 2.6%*p.a.* is $-5.8\% \leq \pi^m \leq 11.0\%$ *p.a.*. We will later see that this range is systematically biased: it underestimates the real equity premium. For now we note that in a review of 150 textbooks discussing the equity premium, Fernández finds estimates for π^m in a range from 3%*p.a.* to 10%*p.a.*, broadly consistent with our predictions [18]. We hasten to add that the 150 textbooks are not consistent in their definitions of the equity premium, so these numbers can only be a rough guide. The simplicity of our input is important: stability and the standard [GBM](#) model for price dynamics.

The equity premium puzzle is an interesting case study of economic science in operation. The literature on the puzzle is large and often takes a psychological and individual-specific perspective. For instance, a more risk-averse individual will demand a higher equity premium. Models of human behaviour enter both into the definition of the equity premium – which lacks consensus [18] – and into its analysis. The problem is treated in wordy essays, and it's also treated using impenetrably complicated formal models. But this additional complexity (compared to our treatment) doesn't add much. It seems to be a result of science incrementally going in one direction and failing to retrace its steps and try a different angle.

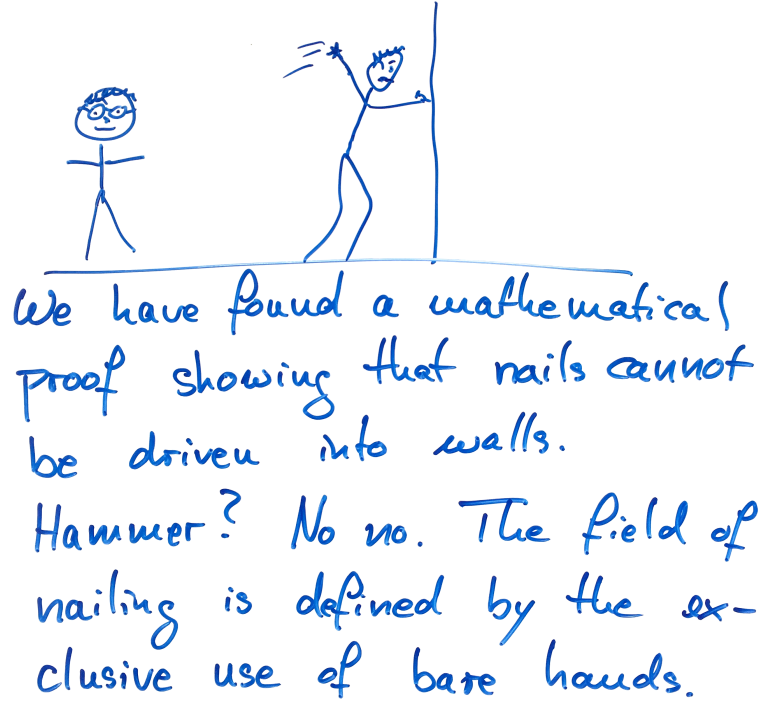


Figure 8.4: The equity premium puzzle.^{fig:cartoon}

The framework we have developed here takes a psychologically naïve perspective, without reference to human behaviour or consumption. It is analytically embarrassingly simple. Only our one key insight was needed: performance over time is not performance of the ensemble. We regard the correct prediction of the equity premium by stochastic market efficiency as the long-sought resolution of the equity premium and price volatility puzzles. It is worth noting that the experts disagree with us: the equity premium puzzle, in their minds, is the fact that observed equity premia cannot be explained with consumption-based asset pricing models. In our view that simply means it's time to throw out consumption-based asset pricing models, whereas clinging on to them can be summarized in the cartoon in Fig. 8.4.

8.3.3 Central-bank interest rates

{section:central-bank}

Our observations are also relevant to the issue of setting a central bank's lending rate. The rate setter would view the total drift μ_s of an appropriate asset or index as given, and the risk-free drift μ_r as the central bank's rate. If the aim is to achieve full investment in productive activity without fuelling an asset bubble, then this rate should be set so that $\ell_{\text{opt}} = 1$. Since $\ell_{\text{opt}} = \mu_s / (\sigma_s)^2$ and $\mu_e = \mu_r - \mu_s$, this is achieved by setting

$$\mu_r = \mu_s - (\sigma_s)^2. \quad (8.30)$$

Using values for μ_s and σ_s estimated in section Sec. 8.4, the optimal interest rate between 1988 and 2018 for the US comes out as 9% *p.a.*, suggesting that actual interest rates over the analysed period have been in a region where they can inflate leverage bubbles (the geometric average on 10-year maturity Treasury bills was about 4.7% *p.a.* lower than this value). The effect is mitigated by

trading costs, meaning that 9% is an overestimate of the optimal interest rate, nonetheless it is hard to deny that asset bubbles were inflated over this period. The task of the central banker can be seen as the task of estimating μ_r and σ_s in the relevant way. This will involve choices about data and timescales which are far from trivial. For instance, in our data analyses at any given time there is an estimate for μ_s and one for σ_s for each possible length of lookback window. Operational matters aside, stability with respect to leverage is an important consideration for any central bank. Leverage efficiency provides a simple quantitative basis for a rate setting protocol and may frame qualitative discussions about interest rates in a useful way.

8.3.4 Fraud detection

{section:fraud}

We will see in the next section that stochastic market efficiency is real. Prices of real traded assets more or less obey the constraints dictated by systemic stability. Of course, over short time scales fluctuations will be large, there are all manner of difficult-to-model effects in real trading environments. These include transaction costs and other operational costs. Still – a realistic asset should more or less satisfy stochastic market efficiency.

But what if it doesn't? What if there is an asset that consistently, and at low volatility outperforms bank deposits? First of all, we now know that such behavior is a challenge to systemic stability. Because of this, we don't expect it to survive for long. Where it does survive for long, one possible explanation is fraud. A common type of fraud is the so-called Ponzi scheme: the fraudster invents an attractive-looking past performance of a portfolio or asset value and entices new investors to give him their money. Instead of investing the money in the non-existent asset, it is passed on to pay off earlier investors and the fraudster siphons off a handy profit. Since the performance of the asset price is invented, there is no strong reason for it to obey stochastic market efficiency. Below we will test the Bernie-Madoff ponzi scheme and ask: was it too good to be true from the perspective of stochastic market efficiency? We will ask the same question about bitcoin: is it too good to be true from this perspective?

8.4 Real markets – stochastic market efficiency in action

{section:real}

In Sec. 8.3 we went through some earth-shattering consequences of leverage efficiency. Of course we wouldn't have done that if there wasn't strong evidence of this organizing principle actually working in practice. Let's get some data and try it out.

We test the stochastic efficiency hypothesis in real markets, quite naïvely, by backtesting leveraged investments. The simplest thing to check is this: take two time series, one of the returns of a very stable, low-volatility, asset, like bank deposits, and the other of the returns of a more volatile asset, like shares. Then try out all possible fixed-leverage investment strategies and see which one does best. The leverage where that happens is a real-world estimate of ℓ_{opt} , which we'll call $\ell_{\text{opt}}^{\text{est}}$.

Of course we can argue about the data until the cows come home – which precise data set should be used, what biases will it introduce, and so on. But

let's set those worries aside for the moment and just try something.

8.4.1 Backtests of constant-leverage investments

An investment of constant leverage is backtested using two time series, $r_s(t)$ (risky returns) and $r_r(t)$ (returns on bank deposits) as illustrated in Fig. 8.1.

1. We start with equity of $x_\ell(t_0; \ell) = \$1$, consisting of ℓ in the risky asset and $\$(1 - \ell)$ in bank deposits.
2. Each day the values of these holdings are updated according to the historical returns, so that

$$\ell x_\ell(t; \ell) \rightarrow r_s(t) \ell x_\ell(t; \ell) \quad (8.31)$$

$$(1 - \ell) x_\ell(t; \ell) \rightarrow r_r(t) (1 - \ell) x_\ell(t; \ell). \quad (8.32)$$

3. The portfolio is then rebalanced, *i.e.* some risky holdings are “bought” or “sold” (swapped for cash) so that the ratio between risky holdings and equity remains ℓ .

In step 1, we could include extra fees for borrowing cash or stock, and in step 2 we could include transaction costs. For the sake of simplicity we leave out these effects and refer the reader to [?]. The result of such extra complexity is two-fold: portfolios with leverage $\ell < 0$ or $\ell > 1$ are penalized by borrowing costs, and any portfolios except those with $\ell = 0$ or $\ell = 1$ are penalized by trading costs. Both effects reinforce stochastic market efficiency, in the sense that they bring $\ell_{\text{opt}} \rightarrow 1$.

But in the simplest instance, steps 1–3 result in the following protocol, which is a discretized version of (Eq. 8.8)

$$x_\ell(t) = \underbrace{r_s(t) \ell x_\ell(t - \Delta t)}_{\text{new risky holdings}} + \underbrace{r_r(t) (1 - \ell) x_\ell(t - \Delta t)}_{\text{new bank deposits}}, \quad (8.33)$$

which we repeat until the final day of the backtest, t_{max} , when the final equity $x_\ell(t_{\text{max}})$ is recorded. Figure 8.5 shows a few trajectories of leveraged investments in the [S&P500 total return index \(S&P500TR\)](#) at different leverages ℓ .

The optimal leverage is that which maximizes the final equity – the trajectories $x_\ell(t)$ are turned into curves $x_\ell(t_{\text{max}})$ by fixing $t = t_{\text{max}}$, see Fig. 8.5. If at any time the equity falls below zero, the investment is declared bankrupt and the backtest is considered invalid for the corresponding leverage. This happens when the leverage is very high and the risky asset drops in value, or when it's very negative and the risky asset rises in value. Between these extremes, a smooth curve emerges, resembling a Gaussian bell-shape (or an upside-down parabola on logarithmic vertical scales).

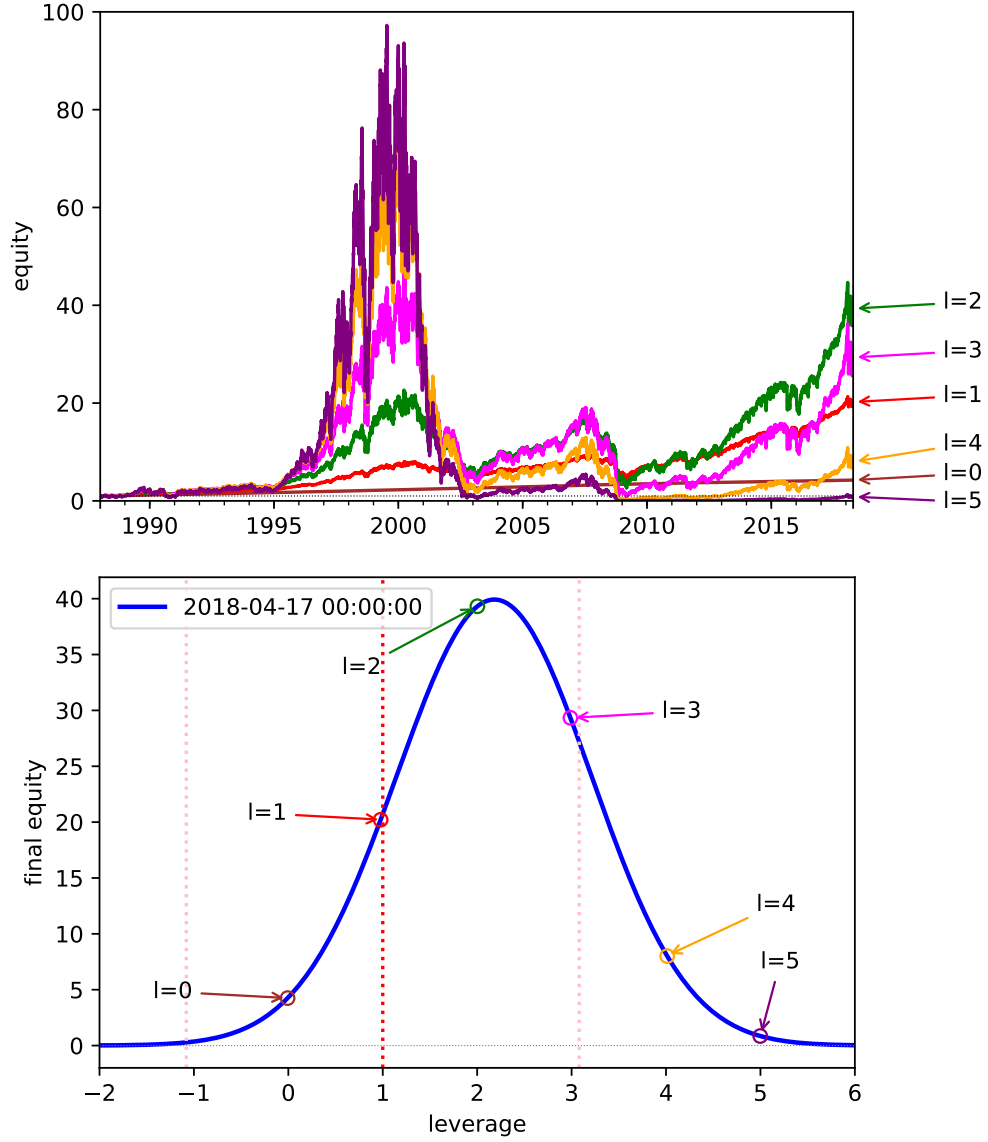


Figure 8.5: **Top:** equity $x_\ell(t)$ as it evolves over time, in investments of initially \$1 in the [S&P500TR](#) at different constant leverages, rebalanced daily as in Fig. 8.1. For riskless returns we use interest rates of 10-year T-bills. Initially, as leverage increases $\ell = 0 \rightarrow 1 \rightarrow 2$ the benefits of investing in the [S&P500TR](#) manifest themselves. But eventually, as $\ell = 3 \rightarrow 4 \rightarrow 5$, harmful fluctuations become so large that returns diminish. Each leverage produces one value for the final equity. **Bottom:** Final equity, $x_\ell(t_{\max})$, for 500 different leverages (where $t_{\max}$ is 1 March 2018 in this case). Each leverage produces a ragged time series of equity (top panel). But the final equity as a function of leverage is a smooth curve. The vertical dotted lines indicate our prediction of the position of the optimum at $\ell = 1$, and the range up to 2 standard errors away. fig:STR_trajectories

8.4.2 Data sets

We want to compare the true performance of one asset to that of another, and therefore it's important to find time series that really resemble changes in equity over time that an investment would have experienced. A stock market index would be nice because it will give a broad overall impression. It won't depend too much on the specific fortunes of any given company.

Strangely, most stock indexes are simply weighted averages of (free-float) market capitalization of a bunch of listed companies, but when one of those companies pays a dividend of \$10, say, and its shares correspondingly drop by \$10, the index drops too. We don't want that, so we'll have to find an index that properly accounts for dividends.

S&P500TR The [S&P500TR](#) includes dividends as desired, so we'll use it as the risky asset and play it against 10-year T-bills as the riskfree asset.

DAX Another index that includes dividends is the German [DAX](#), and we'll try that out too. In this case we will use German overnight deposits for the riskfree asset.

Bitcoin As any scientific model, stochastic market efficiency applies within some range of conditions. We will use Bitcoin to discover the limits of this range. We developed our hypothesis using our knowledge and intuition of economies, stock markets, and central banks. Let's take it outside of this context and see what happens if we use Bitcoin as the risky asset. It doesn't pay any dividends, so we don't have to worry about that, but its performance has been very different from that of a stock market, so it will be informative to see what will happen to our hypothesis.

Madoff We take the idea of finding the limits of the range of applicability one step further, by using an asset that has officially been declared bullshit: investments in Bernie Madoff's Ponzi scheme. Monthly returns were included in a now famous complaint filed to the [Securities and Exchange Commission \(SEC\)](#) by Harry Markopolos in 2005 [30]. The complaint was ignored by the SEC, and Madoff happily continued his scheme until the turbulent financial events of 2008 brought to light what he'd been up to. We digitized the returns in Markopolos's complaint and ran the same backtests for Madoff as for the other assets.

For all assets, we will overestimate real optimal leverage in the most straightforward backtest. Here are some reasons:

1. Whenever trading occurs to rebalance the portfolio, trading costs would be incurred in reality.
2. Money borrowed for leveraging up (and shares borrowed for shorting) come at a premium that the investor has to pay.
3. Stock indexes reflect well diversified portfolios of successful companies – they are affected by survivorship bias because unsuccessful companies leave the index.

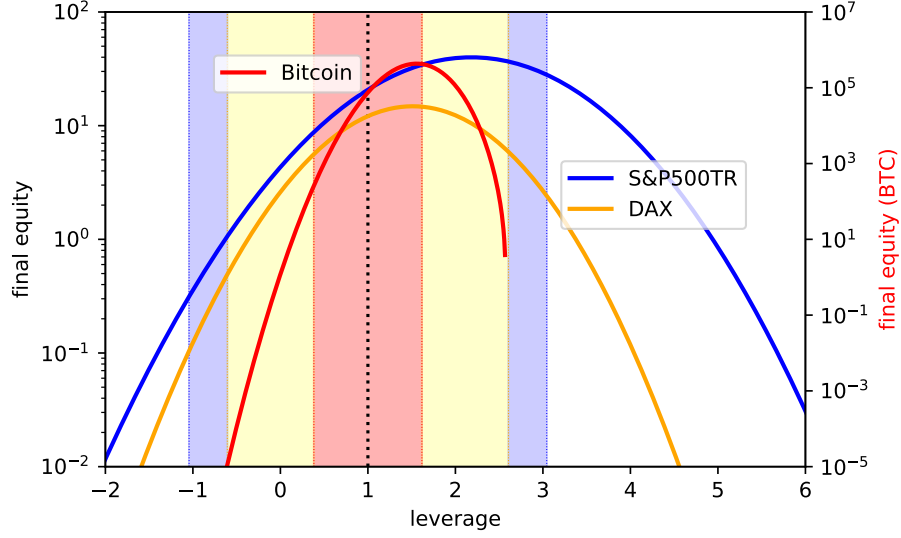


Figure 8.6: Final equity $x_\ell(t_0; \ell)$ for investments of initially \$1 in the S&P500TR (blue), the DAX (yellow), and Bitcoin (red, scale on the right). The maxima of all three curves are near the predicted optimum of $\ell_{\text{opt}} = 1$, where by “near” we mean that they are all within 2 standard errors of the prediction, without any correction for the biases introduced by neglecting trading costs. fig:STR_final_all

We could try to correct for this bias, and we have done that in [44]. The necessary assumptions make the results more realistic but also more dependent on subjective guesses, so we’ll just stick to the simplest case here and note that it will produce an over-estimate of ℓ_{opt} .

8.4.3 Full time series

The bottom panel of Fig. 8.5 shows the leverage-return curve for the S&P500TR. The optimal leverage is a little greater than 1, but still within 2 standard errors from the prediction. Without further ado, let’s add to this figure the same curves for the DAX and Bitcoin, this time on logarithmic vertical scales to see a bit more structure, Fig. 8.6.

The overall impression that emerges is this: leverage efficiency is reasonably satisfied for these assets. The shaded areas in Fig. 8.6 show the predicted ranges, 2 statistical standard errors to the left and right from $\ell_{\text{opt}} = 1$, and all observed maxima of the curves are within these bands. We also see a convincing bias: all estimates are greater than 1, as expected given the bias in the estimate that we mentioned earlier.

From a scientific point of view it’s not very interesting to show lots of cases where a theory gives good predictions. We want to know the limits of the theory. When does it not apply? Where does it break down? Before we looked at the data we thought that Bitcoin might be an example of an asset that doesn’t obey stochastic market efficiency. But it looks just as predicted – unspectacular. We didn’t have a specific reason to suspect that the theory wouldn’t apply to Bitcoin, it just seemed like an extreme asset, unlike the assets we had in mind

when we developed the theory.

But let's see what happens where the theory really doesn't apply. In Fig. 8.7 we add to the curves of Fig. 8.6 the return-vs.-leverage curve for Bernie Madoff's fund. This is instructive because it shows just how different the results could have been. If it weren't for leverage efficiency optimal leverage for the S&P500TR, for example, could be miles away from our prediction. We have to show the Madoff data in a separate figure because the scales are so different that the differences between Bitcoin, S&P500TR, and the DAX become barely visible, and we thought they may be of interest. Apart from showing the limits of the range of applicability of our theory, the Madoff data suggest that leverage-return curves can help detect investments that are too good to be true, *i.e.* fraud.

Of course, Madoff could have chosen the invented returns of his fund in a way that's consistent with stochastic market efficiency, just to have something credible to show to prospective victims of his fraud. But apparently he didn't bother to do that, and people believed that he'd really found a way to print money. As a former CEO of Nasdaq he himself had sufficient credibility, never mind the data.

That stochastic market efficiency really doesn't apply to Madoff's returns is a good sign – showing that a theory fails to predict things it shouldn't be able to predict is important. A theory about asset price fluctuations that applies even to fictitious prices like Madoff's would be suspicious – it would probably not say much about prices at all but just be a statement that's generally true for a much broader class of phenomena. Having said this, introducing realistic transaction costs, leveraging something by a factor of 100 or 200 is extremely expensive, and the Madoff optimal leverage comes down significantly in a more realistic setup. Nonetheless, in Fig. 8.7 Madoff's returns have a statistical signature that is unequivocally different from prices arising through trading that we have investigated.

8.4.4 Shorter time scales

In Sec. 8.2.5 we derived an expression, (Eq. 8.24), for the variance of optimal leverage that we expect to observe in a finite time series, and we used this to arrive at a statistical error estimate, a scale on which to judge whether leverage efficiency holds or not. But we can use (Eq. 8.24) as a prediction in itself: is the sample variance of ℓ_{opt} , measured in many windows of size Δt well described by (Eq. 8.24)? We can't make the time series longer than they are, but we can make them shorter. In Fig. 8.8 we collect statistics of $\ell_{\text{opt}}^{\text{est}}$ estimated in non-overlapping time windows and find good agreement with the model-specific prediction.

The structure we have identified – leverage efficiency – operates on all time scales from weeks to decades. This broad range of applicability supports our theory of noise in Sec. 8.3.1: it is a requirement of stability that asset prices fluctuate, and the shorter the time scales the less significant any economic information becomes. Thinking of price changes as a reflection of new information may be relevant on time scales of years or decades, where the real economic fortunes of companies become apparent. On time scales of days or minutes, let alone milliseconds, price movements usually have nothing to do with economic information.

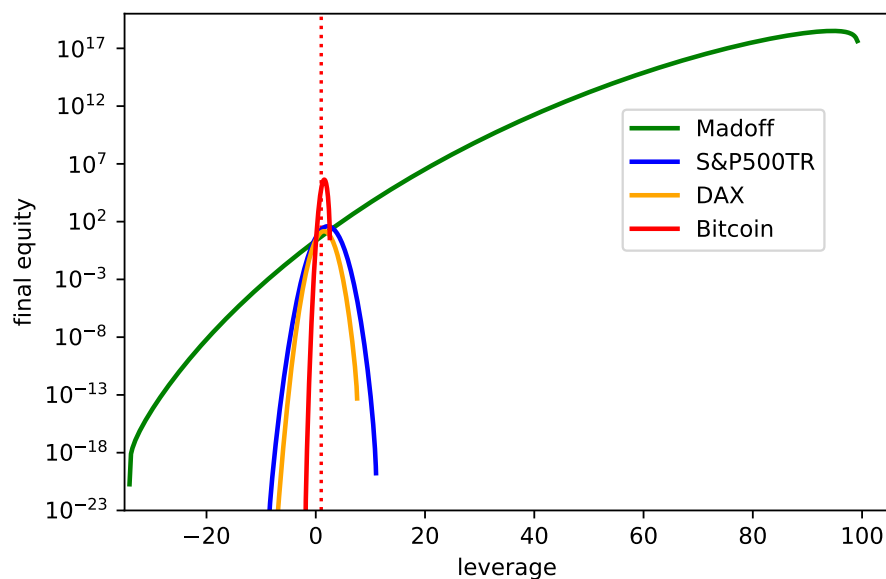


Figure 8.7: Final equity for investments of initially \$1 in the S&P500TR (blue), the DAX (yellow), Bitcoin (red), and Bernie Madoff's ponzi scheme. The maxima for the real assets cluster near the predicted value $\ell_{\text{opt}} = 1$. Madoff's freely invented performance had such low volatility that a leverage of 100 would have been optimal, even if rebalancing could have only occurred once per month. At higher rebalancing frequencies (had the asset actually existed) the green curve would be seen to be the left end of a parabola whose maximum may well be at 200 or 300. fig:STR_final_all_MAD

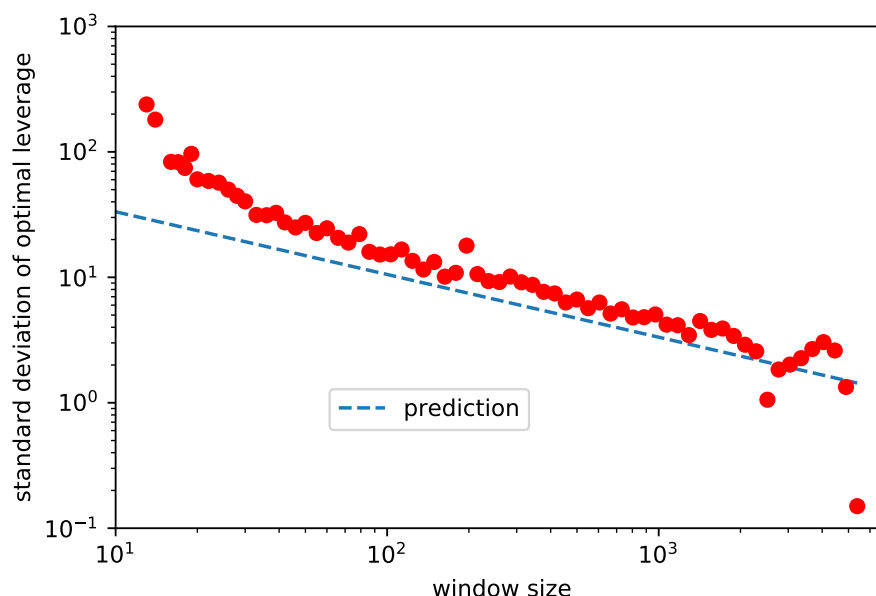


Figure 8.8: The red dots show the standard deviation of samples of $\ell_{\text{opt}}^{\text{est}}(\Delta t)$ as a function of window length measured in days. The blue dashed line shows the prediction of (Eq. 8.24) using the estimate of σ_s obtained by fitting a parabola to the S&P500TR curve in Fig. 8.6. fig:loglog

8.5 Discussion

The primary aim of this final lecture was to demonstrate that the mathematical formalism we have developed to conceptualise randomness in economics, which started with a simple model of wealth evolution by repetition of a gamble, is very powerful. It does more than simply create a rigorous and plausible foundation for economic theory. In particular, because the framework is epistemologically sound, we can make testable predictions, which we can corroborate empirically using real economic data. We could not have guessed from our simple coin-tossing game that we would end up making a prediction about the fluctuations of stock markets or those of bitcoin. The simple but deep conceptual insight that the ergodicity question has been overlooked ultimately shed a good deal of light on the equity premium puzzle, informed the debate of how to set central-bank interest rates, and provided a new way of thinking about noise and the stabilizing role of price fluctuations. What other surprising predictions can we make in this formalism? That, dear reader, is a question for you.

Acronyms

BM Brownian motion.

DAX Deutscher Aktienindex (German stock index).

GBM geometric Brownian motion.

GDP gross domestic product.

LHS left-hand side.

LML London Mathematical Laboratory.

PDF probability density function.

RGBM reallocating geometric Brownian motion.

RHS right-hand side.

S&P500TR S&P500 total return index.

SDE stochastic differential equation.

SEC Securities and Exchange Commission.

List of Symbols

A An observable.

a_v Itô coefficient function for ergodicity mapping $v(x)$.

a_x Itô coefficient function for wealth process x .

b_v Itô coefficient function for ergodicity mapping $v(x)$.

b_x Itô coefficient function for wealth process x .

C Replacement cost of a ship.

d Differential operator in Leibniz notation, infinitesimal.

δt A time interval corresponding to the duration of one round of a gamble or, mathematically, the period over which a single realisation of the constituent random variable of a discrete-time stochastic process is generated..

δ Most frequently used to express a difference, for instance δx is a difference between two wealths x . It can be the Kronecker delta function, a function of two arguments with properties $\delta(i, j) = 1$ if $i = j$ and $\delta(i, j) = 0$ otherwise. It can also be the Dirac delta function of one argument, $\int f(x)\delta(x - x_0)dx = f(x_0)$.

Δ Difference operator, for instance Δv is a difference of two values of v , for instance observed at two different times.

Δt A general time interval..

η Langevin noise with the properties $\langle \eta \rangle = 0$ and $\langle \eta(t_1)\eta(t_2) \rangle = \delta(t_1 - t_2)$.

f Generic function.

F Fee to be paid.

$F^{\max}$ Maximum fee a player is willing to pay.

- g Growth rate.
- G Gain from one round trip of the ship.
- g_{est} Parameter value for time-average growth rate in a general stochastic growth process.
- g_{e} Ergodic growth rate for exponential growth.
- g_{a} Ergodic growth rate under additive dynamics, *i.e.* rate of change, $g_{\text{a}}(x; t, \Delta t) = \frac{\Delta x(t)}{\Delta t}$.
- γ Parameter specifying the value of a time-average growth rate. This enables statements like $g_{\text{time}} = \gamma$, *i.e.* the time average growth rate take the value γ .
- $g_{\langle \rangle}$ Exponential growth rate of the expectation value.
- g_{m} Ergodic growth rate for multiplicative dynamics, *i.e.* exponential growth rate, $g_{\text{m}}(x; t, \Delta t) = \frac{\Delta \ln x}{\Delta t}$.
- $\bar{g}$ Time-average ergodic growth rate.
- $\bar{g}_{\text{m}}$ Time-average ergodic growth rate for multiplicative dynamics, $\bar{g}_{\text{m}}(x) = \lim_{\Delta t \rightarrow \infty} \frac{\Delta \ln x}{\Delta t}$.
- i Label for a particular realization of a random variable.
- j Label of a particular outcome.
- J Size of the jackpot.
- k dummy.
- K Number of possible values of a random variable.
- ℓ Leverage.
- L Insured loss.
- ℓ^- Smallest leverage for zero time-average growth rate.
- ℓ_{opt} Optimal leverage.
- $\ell_{\text{opt}}^{\text{est}}$ Estimate of optimal leverage from data.
- ℓ^+ Largest leverage for zero time-average growth rate.
- $\ell^{\pm}$ Leverages for zero time-average growth rate.
- m Index specifying a particular gamble.
- μ Drift term in [BM](#).
- μ_{r} Drift of riskless asset.
- μ_{s} Drift of risky asset.
- μ_{e} Excess drift.

- n n_j is the number of times outcome j is observed in an ensemble.
- N Ensemble size, number of realizations.
- $\mathcal{N}$ Normal distribution, $x \sim \mathcal{N}(\langle p \rangle, \text{var}(p))$ means that the variable p is normally distributed with mean $\langle p \rangle$ and variance $\text{var}(p)$.
- o Little-o notation.
- p Probability, p_j is the probability of observing event j in a realization of a random variable.
- $\mathcal{P}$ Probability density function.
- φ Reallocation rate in RGBM.
- π^m Equity premium in our model.
- q Possible values of Q . We denote by q_i the value Q takes in the i^{th} realization, and by q_j the j^{th} -smallest possible value.
- Q Random variable defining a gamble through additive wealth changes.
- r Random factor whereby wealth changes in one round of a gamble.
- $g_{\langle \rangle}$ Exponential growth rate of ensemble average wealth in a gamble.
- r_{eff} Effective per round factor by which wealth is multiplied over many rounds of a gamble.
- r_{time} Effective per round factor by which wealth is multiplied as the number of rounds of a gamble diverges.
- g_{time} Exponential growth rate of individual wealth in a gamble as the number of rounds diverges.
- r_r Factor whereby wealth deposited in a bank changes.
- r_s Factor whereby wealth invested in stock changes.
- s Dummy variable in an integration.
- S Sharpe ratio.
- σ Magnitude of noise in a Brownian motion.
- σ_s Volatility of risky asset.
- t Time.
- T Number of sequential iterations of a gamble, so that $T\delta t$ is the total duration of a repeated gamble.
- t_0 Specific value of time t , usually the starting time of a gamble..
- t_{max} Final time in a sequence of repeated gambles.
- τ Dummy variable indicating a specific round in a gamble.

u Utility function.

v Stationarity mapping function, so that $v(x)$ has stationary increments.

var Variance.

W Wiener process, $W(t) = \int_0^t dW$ is continuous and $W(t) \sim \mathcal{N}(0, t)$.

x Wealth.

x_0 Value of an investment in the riskless asset.

x_1 Value of an investment in the risky asset.

x Deterministic wealth.

ξ A standard normal variable, $\xi \sim \mathcal{N}(0, 1)$.

x_ℓ Value of an investment in a leveraged portfolio.

Y Random variable that is a function of another random variable, Z .

z Generic value of a random variable.

Z Generic random variable.

Bibliography

- [1] A. Adamou and O. Peters. Dynamics of inequality. *Significance*, 13(3):32–35, 2016.
- [2] J. Aitchison and J. A. C. Brown. *The lognormal distribution*. Cambridge University Press, 1957.
- [3] G. I. Barenblatt. *Scaling*. Cambridge University Press, 2003.
- [4] J. Benhabib, A. Bisin, and S. Zhu. The distribution of wealth and fiscal policy in economies with finitely lived agents. *Econometrica*, 79(1):123–157, 2011.
- [5] Y. Berman and Y. Shapira. Revisiting $r > g$ —the asymptotic dynamics of wealth inequality. *Physica A: Statistical Mechanics and its Applications*, 467:562–572, 2017.
- [6] D. Bernoulli. Specimen Theoriae Novae de Mensura Sortis. Translation “Exposition of a new theory on the measurement of risk” by L. Sommer (1954). *Econometrica*, 22(1):23–36, 1738.
- [7] F. Black. Noise. *The Journal of Finance*, 41(3):528–543, 1986.
- [8] J.-P. Bouchaud. Note on mean-field wealth models and the random energy model. Mimeo.
- [9] J.-P. Bouchaud. On growth-optimal tax rates and the issue of wealth inequalities. <http://arXiv.org/abs/1508.00275>, Aug. 2015.
- [10] J.-P. Bouchaud and M. Mézard. Wealth condensation in a simple model of economy. *Physica A*, 282(4):536–545, 2000.
- [11] T. W. Burkhardt. The random acceleration process in bounded geometries. *Journal of Statistical Mechanics: Theory and Experiment*, 2007(7):P07004, 2007.
- [12] J. Y. Campbell. *Financial decisions and markets, A course in asset pricing*. Princeton University Press, 2017.
- [13] M. Corak. Income inequality, equality of opportunity, and intergenerational mobility. *Journal of Economic Perspectives*, 27(3):79–102, 2013.
- [14] T. M. Cover and J. A. Thomas. *Elements of Information Theory*. John Wiley & Sons, 1991.

- [15] B. Derrida. Random-energy model: Limit of a family of disordered models. *Phys. Rev. Lett.*, 45(2):79–82, July 1980.
- [16] B. Derrida. Random-energy model: An exactly solvable model of disordered systems. *Phys. Rev. B*, 24:2613–2626, Sept. 1981.
- [17] A. Drăgulescu and V. M. Yakovenko. Exponential and power-law probability distributions of wealth and income in the united kingdom and the united states. *Physica A: Statistical Mechanics and its Applications*, 299(1):213–221, 2001.
- [18] P. Fernández. The equity premium in 150 textbooks. Working paper, IESE Business School, University of Navarra, 2009.
- [19] R. P. Feynman. *The Feynman lectures on physics*. Addison-Wesley, 1963.
- [20] J. M. Harrison. *Brownian motion of performance and control*. Cambridge University Press, 2013.
- [21] J. C. Hull. *Options, Futures, and Other Derivatives*. Prentice Hall, 6 edition, 2006.
- [22] K. Itô. Stochastic integral. *Proc. Imperial Acad. Tokyo*, 20:519–524, 1944.
- [23] R. Kaas, M. Goovaerts, J. Dhaene, and M. Denuit. *Modern Actuarial Risk Theory*. Springer, 2 edition, 2008.
- [24] J. L. Kelly, Jr. A new interpretation of information rate. *Bell Sys. Tech. J.*, 35(4):917–926, July 1956.
- [25] P. S. Laplace. *Théorie analytique des probabilités*. Paris, Ve. Courcier, 2 edition, 1814.
- [26] H. A. Latané. Rational decision-making in portfolio management. *J. Fin.*, 14(3):429–430, 1959.
- [27] M. D. Lee and E.-J. Wagenmakers. *Bayesian cognitive modeling*. Cambridge University Press, Cambridge, UK, 2013.
- [28] S. F. LeRoy. excess volatility tests. In S. N. Durlauf and L. E. Blume, editors, *The New Palgrave Dictionary of Economics*, pages 78–80. Springer, 2016.
- [29] K. Liu, N. Lubbers, W. Klein, J. Tobochnik, B. Boghosian, and H. Gould. The effect of growth on equality in models of the economy. *arXiv*, 2013.
- [30] H. M. Markopolos. The world’s largest hedge fund is a fraud. Filed with the Securities and Exchange Commission, Nov. 2005.
- [31] H. Markowitz. Portfolio selection. *J. Fin.*, 1:77–91, Mar. 1952.
- [32] H. M. Markowitz. Investment for the long run: New evidence for an old rule. *J. Fin.*, 31(5):1273–1286, Dec. 1976.
- [33] H. M. Markowitz. *Portfolio Selection*. Blackwell Publishers Inc., second edition, 1991.

- [34] D. Meder, F. Rabe, T. Morville, K. H. Madsen, M. T. Koudahl, R. J. Dolan, H. R. Siebner, and O. J. Hulme. Ergodicity-breaking reveals time optimal economic behavior in humans. *arXiv:1906.04652*, 2019.
- [35] R. Mehra and E. C. Prescott. The equity premium puzzle. *J. Monetary Econ.*, 15:145–161, 1985.
- [36] K. Menger. Das Unsicherheitsmoment in der Wertlehre. *J. Econ.*, 5(4):459–485, 1934.
- [37] P. R. Montmort. *Essay d’analyse sur les jeux de hazard*. Jacque Quillau, Paris. Reprinted by the American Mathematical Society, 2006, 2 edition, 1713.
- [38] H. Morowitz. *Beginnings of cellular life*. Yale University Press, 1992.
- [39] M. E. J. Newman. Power laws, Pareto distributions and Zipf’s law. *Contemp. Phys.*, 46(5):323–351, 2005.
- [40] V. Pareto. *Cours d’Économie Politique*. F. Rouge, Lausanne, Switzerland, 1897.
- [41] O. Peters. Comment on D. Bernoulli (1738).
- [42] O. Peters. Optimal leverage from non-ergodicity. *Quant. Fin.*, 11(11):1593–1602, Nov. 2011.
- [43] O. Peters. The time resolution of the St Petersburg paradox. *Phil. Trans. R. Soc. A*, 369(1956):4913–4931, Dec. 2011.
- [44] O. Peters and A. Adamou. Stochastic market efficiency. *arXiv preprint arXiv:1101.4548*, 2011.
- [45] O. Peters and A. Adamou. Insurance makes wealth grow faster. *arXiv:1507.04655*, July 2015.
- [46] O. Peters and A. Adamou. An evolutionary advantage of cooperation. *arXiv:1506.03414*, May 2018.
- [47] O. Peters and A. Adamou. The sum of log-normal variates in geometric brownian motion. *arXiv:1802.02939*, Feb. 2018.
- [48] O. Peters and M. Gell-Mann. Evaluating gambles using dynamics. *Chaos*, 26(2):23103, 2016.
- [49] O. Peters and W. Klein. Ergodicity breaking in geometric Brownian motion. *Phys. Rev. Lett.*, 110(10):100603, Mar. 2013.
- [50] T. Piketty. *Capital in the Twenty-First Century*. Harvard University Press, Cambridge, MA, 2014.
- [51] S. Redner. Random multiplicative processes: An elementary tutorial. *Am. J. Phys.*, 58(3):267–273, Mar. 1990.
- [52] P. Samuelson. St. Petersburg paradoxes: Defanged, dissected, and historically described. *J. Econ. Lit.*, 15(1):24–55, 1977.

- [53] P. A. Samuelson. What classical and neoclassical monetary theory really was. *The Canadian Journal of Economics*, 1(1):1–15, 1968.
- [54] A. Sen. *On Economic Inequality*. Oxford: Clarendon Press, 1997.
- [55] M. Sewell. History of the efficient market hypothesis. Technical report, University College London, Department of Computer Science, Jan. 2011.
- [56] W. F. Sharpe. Mutual fund performance. *J. Business*, 39(1):119–138, 1966.
- [57] The World Inequality Database. Usa top 10% and 1% and bottom 50% wealth shares, 1913–2014. <http://wid.world/data/>, 2016. Accessed: 12/26/2016.
- [58] H. Theil. *Economics and information theory*. North-Holland Publishing Company, 1967.
- [59] E. O. Thorp. *Beat the dealer*. Knopf Doubleday Publishing Group, 1966.
- [60] G. E. Uhlenbeck and L. S. Ornstein. On the theory of the brownian motion. *Physical Review*, 36(5):823–841, Sept. 1930.
- [61] N. G. van Kampen. *Stochastic Processes in Physics and Chemistry*. Elsevier (Amsterdam), 3rd edition, 1992.
- [62] O. Vasicek. An equilibrium characterization of the term structure. *Journal of Financial Economics*, 5(2):177–188, 1977.
- [63] J. von Neumann and O. Morgenstern. *Theory of Games and Economic Behavior*. Princeton University Press, 1944.