

De Albuquerque, Emmanuel P.

Working Paper

The creation and diffusion of knowledge - an agent based modelling approach

UCD Centre for Economic Research Working Paper Series, No. WP21/13

Provided in Cooperation with:

UCD School of Economics, University College Dublin (UCD)

Suggested Citation: De Albuquerque, Emmanuel P. (2021) : The creation and diffusion of knowledge - an agent based modelling approach, UCD Centre for Economic Research Working Paper Series, No. WP21/13, University College Dublin, UCD Centre for Economic Research, Dublin

This Version is available at:

<https://hdl.handle.net/10419/237591>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.

UCD CENTRE FOR ECONOMIC RESEARCH

WORKING PAPER SERIES

2021

**The Creation and Diffusion of Knowledge
- an Agent Based Modelling Approach**

Emmanuel P de Albuquerque, University College Dublin

WP21/13

May 2021

**UCD SCHOOL OF ECONOMICS
UNIVERSITY COLLEGE DUBLIN
BELFIELD DUBLIN 4**

The Creation and Diffusion of Knowledge

AN AGENT BASED MODELLING APPROACH

EMMANUEL P. DE ALBUQUERQUE¹
University College Dublin
(School of Economics
and Spatial Dynamics Lab)

Abstract

In this paper I propose a novel abstract mechanism for the creation and diffusion of knowledge and use an agent based modelling approach to explore it. The mechanism takes into account the relation between the phenomena that agents attempt to explain and the stocks of knowledge available in a society, be it individually or collectively. I find that the aggregate number of knowledge units in a society increases more slowly, the more naive its inhabitants are. I also find that the proximity between phenomena plays an important role in how often the same knowledge unit can be used. A discussion on agent based models as a means of insight into society is offered.

Keywords: Agent-based modelling, Cognitive distance, Exploitation, Exploration, Innovation, Knowledge creation, Knowledge diffusion, Learning
JEL Classification: B52, C63, D83, O33

1 Introduction

One of the main contributions that economic geography and regional economics provide to the orthodox study of economics is a concern for the spatial aspects of how economic affairs are conducted. One of the insights from this line of inquiry was that being close, in *proximity*, to other economic agents eased economic activity, innovation and, more specific for the purposes of this paper, knowledge sharing.

Upon expanding this discussion, which initially focused on the spatial kind of proximity (i.e. the physical distance between agents), it became clear that other types of proximity also played a role in the sharing of knowledge, that the proximity among agents was important in multiple dimensions, such as the social, the cognitive, or the institutional, among others.

¹Acknowledgements: this paper is an extended version of an essay for the lecture on Agent Based Models at UCD, by professor Pablo Lucas, whose comments on initial drafts were invaluable. The author is a PhD student funded by Science Foundation Ireland (SFI).

In the current paper I develop a formal model, which allows this multi-dimensional aspect of proximity to be studied in relation to knowledge creation and diffusion. To do that, while still being able to analyse group, network and individual dynamics, an agent based modeling (ABM) approach seemed to be ideal. The application of ABMs to problems focused on the interaction of idiosyncratic agents is widespread, with a simple, early example in Wilensky (1998). An important discussion on the external validity of ABMs can also be found in Casini and Manzo (2016).

The current model simulates an environment in which individuals (agents) must deal with phenomena that randomly occur to them. To do that, they may resort to their individual knowledge stock, to a collective knowledge stock or engage in research, creating new knowledge.

Boschma (2005) reviews the role of proximity as a driver of innovation. There, surveying the contributions of others in the literature of economic geography, he expands upon the intuitive notion of proximity as a measure of spatial distance, highlighting that proximity may also be social, cognitive, institutional or organizational.

Schlaile et al. (2018) proposes a discussion regarding the different methods for representing knowledge in modeling attempts. There the authors argue that their approach, i.e. creating a set of uniquely identifiable knowledge units related with one another through a network, is an improvement to modeling knowledge as a scalar, which would represent the stock of knowledge an agent possesses, or as a vector of scalars, which would represent the stock of knowledge in various dimensions or *categories*.

The modeling approach I employ in this proposal is defining *each* knowledge unit as a vector in a metric space,² which yields an uniquely identified unit as in Schlaile et al. (2018), but which also allows for more versatile calculations of distance, other than only network related measures (e.g. degree centrality, shortest path length, etc.), which are nonetheless also possible in the proposed framework.

Moreover, Schlaile et al. (2018) models each knowledge unit as a bit string of length n_k . This means that there are 2^{n_k} possible configurations for one knowledge unit, the distance between each of them being measured by the normalized sum of the XOR operator over the n_k bits. Which means in essence that, in each dimension, in each bit, a knowledge unit can only have a dichotomous relation to another, being either equal or different.

Overall, the previous approaches to modeling the creation and diffusion of

²So that each knowledge unit an individual possesses is a vector and his knowledge stock is the collection of all such vectors. It is useful to compare it with the vector approach of the last paragraph, where knowledge units of a same category are added up together indiscriminately.

knowledge focused on the relationship between knowledge units, relegating the relationship of the knowledge units with the phenomena which motivate them, to the background. However, there are aspects of the creation and diffusion of knowledge which are only possible to investigate by focusing on this relegated relationship.

For instance, by analysing only the compatibility of received knowledge, it is possible to observe how the knowledge stocks of different agents interact with one another, but the questions of how the received knowledge relates to the actual phenomenon which motivated it are left unanswered. This paper is suited to fill this gap in the literature.

More specifically, the present model is a simplified account of how agents create and share knowledge, based upon one overarching mechanism: that agents create and consume knowledge according to its accuracy. It assumes that all the explanations, representations, conceptualizations (...) we may have for the phenomena that surround us in the world carry with them one attribute: accuracy regarding its relation to the phenomenon it aims to explain, to represent.³

It is observable that people create somewhat accurate explanations about the world and those who deem them worthy consume it. Some of us create very accurate explanations - consider the top scholars in your field of research -, some of us don't. Some of us are very demanding when it comes to accepting other people's accounts or explanations, some of us aren't.

To illustrate, consider Newton's gravitational theory of motion: it is a 17th century explanation of how and why celestial bodies move, it produces, however, inaccurate predictions regarding the movement of Mercury (Harper, 2007). Perhaps those who wished to predict the motion of Jupiter deemed it accurate enough, while those who wished to predict the motion of Mercury did not. It is a piece of knowledge created with a certain accuracy and consumed according to this attribute (among other things, surely, such as the reputation of its creator, the dissemination of the idea, how it integrated with previous understandings of the world, etc.).

One interesting result from this proof of concept is that, in the model, the production of more accurate knowledge discourages the formation of new understandings regarding the same phenomenon.

I believe this line of inquiry is as parsimonious as it is useful to anyone who wishes to investigate the history and evolution of ideas. It can shed light into the consequences of the rise of the scientific method: how it not only increased the accuracy of understandings in the world, but also how it changed how demanding we are of the information we receive from others. Moreover, we can inquire into, if the

³ More on how I define phenomena, knowledge and accuracy in the following section, where I specify the model.

scientific method is more amenable to some objects of study than to others - specially if we consider questions pertaining to the subjectivity of existing, to existentialism or spirituality, to exemplify -, how do we compare the evolution of these different bodies of knowledge.

The remainder of this paper is structured as follows: section 2 explores agent based models, with subsections 2.1 and 2.2 displaying self contained discussions regarding (2.1) analytical sociology's relation to ABM models and (2.2) issues of causality and validation; section 3 specifies the model, with a subsection dedicated to the algorithmic workings of the simulation; section 4 discusses the results; and, lastly, section 5 presents the concluding remarks and directions for future research.

2 Agent based models

Agent Based Models (ABMs) are models. This is a truism, of course; but a necessary one. Necessary because (from my perspective), if we wish to understand ABMs, remembering that they are models and what models are in essence is the mandatory first step. A step that, once taken, makes understanding ABMs merely the case of defining what *agent based* stands for. It is a small detour, but one which, I believe, will make matters much clearer by the end of this essay.

The route I will take is to explore my understanding of what models, in essence, are, and then relate them to ABMs. Then, I will introduce some of the concerns related to the usefulness of models in general, again turning to observe how those concerns relate to ABMs. So we begin.

Models are ubiquitous in human life. To reason is to create *representations* and *abstractions* of the objects of our reasoning. To think of our mothers or fathers, to use an example we can all relate to, we must first consider within what framework we will think, what aspects of their existence will be entertained. In essence, we must choose what is important in light of our reasoning purposes. We can think as sons and daughters, focusing in the way we have felt and made feel in the years of relationship. We can think in an evolutionary and genetic perspective, focusing on the receding hairline or height inherited. We can even think in light of indirect stories told about them in their youth. So that, in essence, to think of them, or anything, we must focus our attention, abstracting and representing, creating models in our mind.

To illustrate the point further, in a less sentimental setting, consider Sir Isaac Newton, as legend tells, sitting under a tree, being hit overhead by a murderous apple. As it hit him, he could have thought of how it felt, or might have remembered how hungry he was or even how peculiar the fruit looked. Instead - and this, of course,

is an anecdotal account to help exemplify my point - he abstracted from the tree, from the soil he was sitting on and even from himself, to focus on the relationship between the mass of that apple and the mass of the planet. By representing (i.e. thinking of the apple and the earth as things that possessed mass) and by abstracting (i.e. forgetting all else to focus on those representations and their relationship) he modeled gravity, in one of the most successful models in the history of science.

The models we use in everyday life to think about daily affairs are loosely constructed models, based also on gut feelings and other subjective inputs. Sir Newton’s model is a formal mathematical model, with precise assumptions (i.e. abstractions) and definitions (i.e. representations). And, though in scientific inquiries we try to focus on the more formal models, they are all thought processes, more or less structured, that assist us in reasoning and navigating in the world, be it in everyday life or in astronomical calculations.

Agent based models, as well, are thought processes, but ones which are formally defined, with precise and explicit assumptions and definitions, with a specific focus on agents, on the way these agents *interact* and on the properties that *emerge* from those interactions.

To explore what agent based stands for, in this instance, I will make use of the definition provided by Schlosser (2015), according to whom agents are beings who possess the capacity to act, so that agent based models would be defined models based on beings that act. More than that, ABMs would be focused on those beings that act acting to and with one another, i.e. *interacting*, and also on the properties that emerge from those interactions. The importance of the focus on emerging properties is that action is not performed in a vacuum, always within a context, which often guides them. One aspect of an emerging property is that it relates and influences the context surrounding the agents, the environment in which actions take place.

To illustrate what I mean by a specific focus on agents’ interactions and emergent properties, consider the abstract model for knowledge explored in this paper.⁴ Agents, when learning something new, update their individual and collective knowledge spaces, and, because of this update to the collective knowledge space, they indirectly interact with each other. Through this interaction, a relation between the agents’ parameters is observed, highlighting an emergent property from the model - namely that when the precision of new knowledge creation is higher than the satisfaction requirements of current knowledge, the collective production of knowledge decreases.

⁴As an alternative, consider the flocking model presented in Wilensky (1998), where interacting agents (birds), through simple alterations in their heading (considering parameters of *cohesion*, *alignment* and *separation*) create the common V-Shape formation observed in bird flocking.

Moreover, understand that it is because of the agent based approach to this model that I could focus entirely on the agents' algorithmic decision making and interaction. If we consider, for example, the usual representative agent models from the macroeconomics literature (see Kirman (1992) for a critical discussion), we can see that it would be impossible to use such a framework as a thought process assisting tool to deal even with the simple heterogeneous agents in my knowledge model. I mention this alternative approach to modeling, not because I believe ABM to be superior to it, but to punctuate that, as with all modeling approaches, what is central is what we choose to represent and abstract from.

As I said, ABMs' focus is on agents, interactions and emergent properties, and this focus is important because it allows sociologists, for example, to pair their theoretical accounts, often ridden with implicit assumptions and representations, to another thought process that is amenable to their object of study - namely the individuals in a society, their interactions, the patterns that emerge from their interactions and the way in which these patterns affect the individuals in return -, but one which is formal, precise and explicit.

Another aspect of models in general which is important to highlight is that, because they are thought processes, like these, they need not have any relation to reality; so that, in the same way I can think of a hopping unicorn flying an airplane, I can build a model of it. In a scientific setting, however, and also in sociology, thoughts and models must be related to some aspect of reality so that we may use the insights that they bring to enhance our understanding of the world around us; in a way that a considerable amount of effort surrounding scientific models is in making sure that they are related to reality, in an effort called validation.

And yet another aspect of models in general which is relevant to scientific inquiries is called causality, which, shortly put, is concerned with the relations between representations in the model, specifically if they are causal. For instance, wondering if an increase in the mass of the earth *causes* an increase in the falling acceleration of the apple is considering a causal relation. Causality concerns are particularly important because they allow us to better understand the processes and mechanisms underlying the occurrence of events in our thought process.

To finalize this segment of the discussion, note that validation and causality may come together allowing us to link the insights we make regarding the relations in the model (e.g. are they causal?) with reality (i.e. is the model related enough to reality so that I can extrapolate the model's insights to explain the world). Note also that, these considerations about models, of course, are also pertinent regarding ABMs and, in this essay, I will explore these considerations specifically in relation to ABMs.

In the following section (2.1), we will discuss why ABMs and their focus are important for the analytical sociology approach to sociology, specially in regards to middle range theories; in section 2.2 we will discuss what can be the real world applications of these models with a discussion on validation and causality.

2.1 Relation to analytical sociology

Now, following our discussion in the introduction, we will explore the relationship between AMBs and analytical sociology, specially in regards to middle-range theories. To do that, we will need first to explore a scheme of sociological approaches to understand what a middle-range theory is; then, in a similar fashion, we will explore the possible uses for ABMs; which should allow us to, finally, build a connection between middle-range theories of sociology and middle-range ABMs.

To begin, let us observe the classification of sociological approaches proposed in Hedström and Udehn (2009). There, following a Mertonian view, the authors observed that approaches vary in regards to their generalizability and to their degree of abstraction, forming a two dimensional continuum relating how much they explain (how general) and how much they use to explain (their degree of abstraction). For instance, Sir Newton's gravity theory would be a highly general and highly abstract theory, using little to explain much. It is noteworthy that, in the specific realm of sociology, the authors argue, the relevant opposite of a general theory is an applied, empirical analysis, so that, while gravitational theory is general, the observation of the planets is specific, irregardless of how many explanatory variables are used.

Within this continuum, the authors highlight the role of middle-range theories, which, in loose terms, are fairly abstract accounts that explain a fair deal of sociological phenomena, being conceptualized in distinction from the extremes of the continuum - i.e. the all encompassing, over ambitious theories of everything and the stand alone empirics, be they over generalizing or mere descriptions. In essence, defining them as theories that do not to over emphasize the importance of any one aspect of social behavior (e.g. optimizing behavior of individuals, as in the work of Gary Becker), with contextualized objects of study (i.e. no attempts at constructing laws that explain all social behavior).

Also, one of the most important aspect of the Mertonian approach for our interests in this essay, which is also highlighted in Hedström and Udehn (2009), is their concern with macro to micro links and micro to macro links. The notion, also highlighted in Coleman's boat, that macro structures affect individuals, who affect each other and again affect the macro structure, is an understanding of how societies function which is specially amenable to modeling with an agent based approach.

Squazzoni (2012), in fact, maps the distinctive ABM approaches in a similar fashion to Hedström and Udehn (2009), which, in short, consists of a classification ranging from abstract models on one end to applied, empirically informed models on the other. And there we can see the great extent to which ABMs have been applied to sociology. Also, within this spectrum, we can find middle-range ABMs, which are defined as “empirically grounded theoretical models intended to investigate specific social mechanisms that account for a variety of empirical phenomena that share common features.” (Squazzoni, 2012, p. 24).

It is interesting to note the importance of empiricism in relation to middle-range ABMs, a topic which will be discussed in the next section, under the overarching concept of validation. Also interesting is the highlighted intention of middle-range ABMs to investigate mechanisms, which prompts me to mention the idea of generative causality, also discussed in the next section.

In all, middle-range theories are reasonably ambitious and contextualized theories, which aim not to over stress their explanatory variables, and ABMs are specially suited to these kind of theories for mainly two reasons: first, as mentioned before, ABMs and middle-range theories share an analogous focus on agents, interactions, emergent patterns and the effects of these emergent patterns on agents; second, because ABMs are run through simulations, they allow for problems that do not possess analytical solutions (the majority of them) to be analysed in a systematic way.

Moreover, us humans seem to like to think that we possess some type of control over our behavior, that we possess agency, even if imperfect. This is, obviously, an important aspect of middle-range theories. To analyse how individuals respond to macro changes on a micro level, for instance, we must model the *reaction* of those individuals to those changes; and the same can be said to be true regarding the relations between the individuals.

And, while other formal analytical approaches to modeling struggle to account for the idiosyncrasies of agents and their particular contexts, as must seem quite obvious at this point, *agent* based models are focused precisely on agents and their agency. To briefly illustrate, consider the function *own* in NetLogo, to attribute values to each turtle or each patch. Because of this, ABM is able to follow agents as they diverge into different paths, converge again, etc., and model this salient feature of societies of particular interest to middle-range theories.

2.2 A discussion on causality and validation

Hedström and Udehn (2009) proposes an interesting discussion regarding dif-

ferent conceptualizations one may have of causality, referring to causal laws, statistical associations, counter-factual analysis and causal mechanisms. Here, I will briefly consider this discussion and, then, discuss how ABMs could contribute to it.

Causal laws are general relationships of causality of the type if A then B, valid irregardless of context - a very strong relationship. The law of gravity proposed by Sir Newton, among several others in the natural sciences, have long been used as an example and been associated with the notion of causal laws because of its generality and ubiquity in the *universe*. However, even the classification of these, which are surely the most generalizable of human perceived regularities, as causal laws encounter opposition; the argument being that even them are context dependent (e.g. how does the law of gravity operate in a vacuum? WRONG). Interesting as this philosophical discussion may be, it would escape the scope of this article if not cited merely to briefly illustrate what causality is and how one would proceed to establish it or question it.

A statistical association concept of causality would consist as the persistence of a statistical correlation between phenomena, even after controlling for possible confounding factors. It poses much softer requirements than causal laws, but, because of it, it is more prone to false positives, i.e. interpreting relations as causal when, in fact, the correlation is spurious.

A counter-factual discussion of causality focuses on *what if* considerations of the sort: what would have happened to the stock of cheese had the mice not gotten in? Which involves the factual (i.e. the mice got in), the counter-factual (i.e. the mice did not get in), and the comparison between the two *states*. If with mice and without mice the stock of cheese is unchanged, then there is no causal relationship. If, however, we find that the stock of cheese would have been depleted either way, because, before the mice came in, a hungry Swiss decided to make a fondue, the causal effect of the mice getting in would be null. Or, on the other hand, if we establish that, other than the mice, no one else had access to the cheese, that there were no trap-doors under the fridge, or any other competing hypothesis for its disappearance; so that, had the mice not gotten in, the cheese would have been intact; then we may say that the causal effect of the mice getting in was the disappearance of the cheese. In essence, then, we may say that the counter-factual analysis consists of comparing what would have happened in different states of nature.

A discussion of causal mechanisms, then, would dive deeper and attempt to establish what could have happened in the specific context to warrant a causal claim, focusing on how the relation took place. In the case of the mice above, for instance, even if there were no competing hypothesis as to how the cheese disappeared, an analysis of causal mechanism then would attempt to explain how it would be that

the mice made the cheese disappear. One possible mechanism of disappearance in this case being: the mice ate the cheese.

Now, turning to analyse how ABMs relate to this discussion, note that in order to propose the mechanism above, we needed some parting ground as to propose that the mice ate the cheese, namely a theory or model relating them. After all, returning to the discussion from the introduction, it is not because it is common sense that mice eat cheese, that it would be any less of a thought process, with implicit and explicit assumptions about it, and thus, possibly amenable to formal modeling. ABMs, then, as before, could be the pertinent choice of approach whenever the core features of the mechanism being proposed are focused on agents, their interactions, emergent properties and their feedback.

Given that we settle on some notion of causality, we can assess if the relations being addressed within the ABM are causal. If they are, we may say that there is internal validity of the relations in the model, i.e. that within the bounds of its assumptions and definitions the causal relations will hold, will be valid.

Another very important contribution of ABMs to the analysis of causality is its ability to explore the notion of *consequential manipulation*, which consists of examining causality by manipulating proposed causes and observing the resulting consequences (Casini and Manzo, 2016), an analysis which ABM, as a formal simulation, is perfectly suited to perform.

Casini and Manzo (2016) also refers to the definition of causality as a *generative process*, which consists of examining causality by observing what outcomes could be produced by a model configured in a certain way. To illustrate I can say, in light of my knowledge model, for instance, that under its configurations, in no circumstances, could there be an agent who knew everything, while the rest knew nothing. It would be possible to, changing the assumptions, examine what would be necessary for this scenario to emerge. However, as is, my model could not inform me of a generative process for this scenario. On the other hand, it can inform me of a process in which the amount of collective knowledge changes in respect to differences between the satisfaction and precision parameters. In essence than, I can say that causality as generative processes is concerned with what can be generated, what could be the outcome of causes (or assumptions in a model). I will refer to generative processes once again, at the end of the section.

The next concern related to our scientific explanations are related to validation. This concern can be illustrated in the following manner: consider, for instance, in light of what we just discussed surrounding the application of ABMs to provide accounts of mechanisms, what would assure us that the formal, explicit, agent focused account provided by the ABM to explain the mice eating cheese has anything

to do with reality? The discussion surrounding validation, then, can be said to be a discussion about how valid a link between the models we build and reality is. To once again, refer to the discussion above, we may say that a model well linked to reality is an externally valid model or possess external validity.

To the purposes of this paper, following Casini and Manzo (2016), we will discuss three broad ways one can proceed in order to validate an ABM, namely theoretical realism, empirical calibration and empirical validation.

Theoretical realism consists in basing the representations and abstractions in your model on well established theoretical discussions in the literature. For instance, if I were to model the behavior of an agent as rational, utility maximizing I could use the literature in economics to inform my modeling choices.

Empirical calibration consists of defining the domain within which the parameters of your model can vary according to measures obtained in empirical studies. Imagine, for instance, that in my mice model, I wish that each mouse be restrained in respect to the amount of cheese it can ingest, I can inspect the literature on the eating habits of mice to calibrate my model.

And empirical validation, finally, consists in comparing the results of my model with the pertinent statistics. To illustrate, consider the predictions from Sir Newton's model regarding the position of planets in a given time. A way to empirically validate his results would be to use a telescope to see if the modeled positions matched the actual positions. This example highlights that the validation process need not be perfect for a model to be accepted and useful, as there was sensible discrepancy between the predicted and actual orbit from Mercury, and Newtonian mechanics remain useful to this day.

In all, as a conclusion to this discussion, we can see that ABMs can be specifically useful as an empirically or theoretically validated way to construct theories of mechanism. Consider that, because a mechanism consists of a relation between two states, one previous to the causal effect and one posterior, the modeler is able to validate the inputs of his model with empirical calibration, base the assumptions of relationship in a theoretically realistic way and, finally, empirically validate the model's results against aggregate data.

Particularly pertinent to this discussion on the use of theoretically realistic and empirically calibrated ABMs, is the notion of causality as a generative process referred to above. When properly calibrated and well oriented and contextualized, ABMs can be used to generate predictions, examine bandwidths of likelihood of events, and also examine the probability that the causes of an observed event, in a specific mechanism proposed in the model, are X or Y. For instance, if a theory proposes that a single mouse is eating 5kg of cheese every night, I can empirically

calibrate the model to account for how much a mouse eats and see that it would be impossible. Alternatively, I could want to examine how much cheese 10 mice could eat, given that they fight among them. And if the calibration and theoretical foundations are proper, I should get plausible, informative responses.

So that, if we are focusing on some specific relation between two variables being analysed by the model, identifying the mechanism through which one affects the other, and exploring in which instances the relation holds is investigating the internal validity of the relationship. Linking this relationship with reality through the efforts of validation discussed above is establishing the external validity of the model. To bring back the discussion to middle-range theories, we may also say that all middle-range models can only be externally valid within the specific context to which the analysis is aimed.

3 The Model

3.1 Intuition and simulation algorithm

To model this accuracy mechanism, we assumed, quite reasonably we expect, that there are things for agents to explain, that every agent may know some things and that every agent may be part of a group of agents who share knowledge with each other. To model these assumptions, we define an Euclidean metric space called the *knowledge space*, where three distinct sets interact with each other, namely: the *phenomena set*, the *individual knowledge stock* and the *collective knowledge stock*. The phenomena set (P) is an n -dimensional collection of *phenomena*, with each phenomenon being a n -tuple - these are the things agents attempt to explain. In turn, the individual knowledge stock (KI) and the collective knowledge stock (KC) are both n -dimensional collections of *understandings* - these are the agent's attempts at explaining.⁵

Now, if we go through the fairly straightforward algorithmic workings of the model, I believe the way in which the modeled spaces relate to each other will become clear. The model functions as follows:

1. an agent is prompted with some aspect of reality;
2. she examines her individual knowledge stock to see if she can explain it to her own satisfaction; if yes, the run stops and another agent is prompted;

⁵Lieto et al. (2017) reviews approaches in dealing with cognitive architectures, endorsing *conceptual spaces*, a framework proposed by Gärdenfors (2004), which is similar to our approach.

3. if not, she consults those who are in her group, examining the collective knowledge stock, to see if they can explain it, also to her satisfaction; if yes, she learns this explanation, the run stops and another agent is prompted;
4. if not, she comes up with a new explanation of a certain accuracy. She then learns this explanation, which is now available to her group, the run stops and another agent is prompted.

To conclude the conceptual explanation of the model, the only remaining issues are: i) what determines if an agent is satisfied or not?; ii) how is the accuracy of a new explanation defined (as in step 4 of the algorithm)?

Satisfaction is determined by analysing the relation between the Euclidean distance between a phenomenon and an understanding and a threshold parameter of satisfaction. The accuracy of a new explanation is determined by an accuracy parameter defining the radius of a uniform distribution centered in the phenomenon of interest. In the next subsection, we propose a more precise, mathematical description of the model.

3.2 Mathematical description

Consider an environment populated with a finite set of agents $A = \{1, 2, \dots, N\}$, with the generic agent being denoted by a . Time is discrete with the generic time-period⁶ denoted by $t = 1, 2, \dots$. Consider a finite set of collectives $C = \{c_1, c_2, \dots, c_C\}$, where $c_c = \{1, 3, a\}$ if agents 1, 3 and a are members of the collective c .

Consider a knowledge space, represented by a n -dimensional euclidean space with arbitrary bounds, populated with the *phenomena set*, a set of uniformly distributed n -dimensional vectors, representing the phenomena of reality. Each vector in the phenomena set, denoted P , is called a phenomenon, denoted $phen_{at}$ - to highlight that each phenomenon occurs to some agent a at some time t -, characterizing the set as $P = \{phen_{at} : phen_{at} \in \mathbb{R}^n\}$.

Each agent is characterized by her individual and collective knowledge stocks, denoted KI_a and KC_c , respectively. The individual knowledge stock is formed when the agent learns, either by creating new knowledge or by learning from the group.

New knowledge is created by a function called $creation(phen_{at}) : \mathbb{R}^n \rightarrow \mathbb{R}^n$ such that:

$$k_{at} = creation(phen_{at}) = phen_{at} + \epsilon_r \quad (1)$$

⁶We use time-period and tick interchangeably across the paper.

where k is a n -dimensional knowledge vector, ϵ_r is a noise vector uniformly distributed over a n -ball of radius r centered on $\mathbf{0}$. In other words, the agents create new knowledge by appending a noisy version of the phenomena they encounter, and the smaller the radius r , the smaller the expected noise.

Once a new knowledge is created, it is appended both to the individual knowledge stock of the agent who created it and to the collectives to which she belongs to. So that an agent knows $KI_a = \{k : k \in \mathbb{R}^n\}$, and anything an agent knows is available to her collective, i.e. $KC_c = \bigcup_{a \in c} KI_a$.

Thus, learning from the collective becomes possible, because, since the other agents from the group may have been confronted with different phenomena, the union of all individual knowledge stocks can be larger than a single individual knowledge stock. Learning from the collective means that an agent a appends the knowledge of interest from KC_c to KI_a .

When confronted with a phenomenon, to decide whether to create new knowledge, to use the individual knowledge stock or to learn from the collective, an agent follows an **algorithmic procedure**: **first** she observes her individual knowledge stock and finds the knowledge unit which is closest to the phenomenon, which is given by:

$$k_a^* = \underset{k}{\operatorname{argmin}} d(k, phen_{at}), \forall k \in KI_{at} \quad (2)$$

where k_a^* represents the knowledge vector that minimizes the euclidean distance function $d(k, phen_{at})$ over the individual knowledge stock.

But, since it is possible that the closest knowledge unit is still too far away, the agent compares it to s , an exogenously defined scalar called *satisfaction parameter*, with two possible outcomes: k_a^* is satisfactory or not. More specifically, I define an indicator function $I_{at}(d(k, phen_{at}), s) : \mathbb{R}^2 \rightarrow \{0, 1\}$, such that:

$$I_{at} = I(d(k_a^*, phen_{at}), s) = \begin{cases} 0, & \text{if } d_a^* > s \\ 1, & \text{if } d_a^* \leq s \end{cases} \quad (3)$$

i.e. if the distance from the closest knowledge unit k_a^* to the phenomenon $phen_{at}$ is smaller than the satisfaction parameter s , k_a^* is deemed satisfactory and $I_{at} = 1$. If the distance is bigger however, with $d(k_a^*, phen_{at}) > s$, the indicator function returns zero, with $I_{at} = 0$, meaning that even the closest knowledge unit in KI_a is still unsatisfactorily far from $phen_{at}$.

Then, if $I_{at} = 0$, the agent does the same with her collective knowledge stock, which would lead us to find:

$$k_c^* = \underset{k}{\operatorname{argmin}} d(k, phen_{at}), \forall k \in KC_{ct} \quad (4)$$

where k_c^* also represents the knowledge vector that minimizes the same euclidean distance function $d(k, phen_{at})$, but over the collective knowledge stock KC_{ct} ; and to define the indicator function $I_{ct}(d(k, phen_{at}), s) : \mathbb{R}^2 \rightarrow \{0, 1\}$ given by:

$$I_{ct} = I(d(k_c^*, phen_{at}), s) = \begin{cases} 0, & \text{if } d_c^* > s \\ 1, & \text{if } d_c^* \leq s \end{cases} \quad (5)$$

which similarly to before, indicates whether there are any previous knowledge vectors in the collective knowledge stock KC_{ct} that would be satisfactorily used as an understanding to the phenomenon $phen_{at}$. If $I_{ct} = 1$, there is some knowledge unit in which the satisfaction parameter s is satisfied.

In the alternative case, in which $I_{ct} = 0$, the agent creates a new knowledge as specified by the $creation(phen_{at})$ function.

In summary, then, the algorithmic procedure described above consists of an agent first looking inside his own knowledge set, then, if not satisfied, looking into the collective knowledge set and, if still not satisfied, creating a new knowledge unit with expected noise given by the precision parameter r .

Within this procedure, it is possible to specify the individual knowledge stock $KI_{at} \forall i \in A$, where A is the number of agents, as

$$KI_{at} = \{k : k = k_c^* I_{ct} (1 - I_{at}) + (1 - I_{at}) (1 - I_{ct}) creation(phen_{at}), \forall t \in \{0, \dots, t\}\} \quad (6)$$

which highlights that the individual knowledge stock is made up of knowledge vectors that may come from two different sources: if $I_{at} = 0$ and $I_{ct} = 1$ there is no satisfactory knowledge vector in KI_{at} , but there is in KC_{at} , meaning that the agent learns from the collective; and, if $I_{at} = 0$ and $I_{ct} = 0$, the agent creates a new random knowledge vector which is a noisy version of $phen_{at}$, with a maximal distance r in each dimension.⁷

This is the basic framework of the model.⁸ In the next section I will produce

⁷In the remaining case, in which $I_{at} = 1$, i.e. when there is already a knowledge vector in KI_{at} that satisfies agent a , there is no alteration in the composition of the individual knowledge stock.

⁸The algorithmic structure of the model is liable to discussion, so further developments of the model should attempt to relax it. For instance, it is not clear to me that agents look in the collective stock before attempting to create, nor it is clear to me that agents, which are satisfied with their individual knowledge stock, do not attempt to create even more precise knowledge or find one in the collective. This latter case could be modeled as a tick within a tick, where the agent attempts a more precise knowledge and comes back to the individual knowledge stock in case of failure (an expanding s for each *subtick*).

some initial results, discuss its properties and explore some possible caveats to the analysis.

4 Discussion

4.1 Quantity of knowledge

To begin the analysis of the model I will define a cardinality function $\#(KI_{at}) : KI_{at} \rightarrow \mathbb{N}$ which tells us the number of elements in KI_{at} , with

$$\#(KI_{at}) = \sum_t (1 - I_{at})I_{ct} + (1 - I_{at})(1 - I_{ct}) \quad (7)$$

Given the stochastic nature of the model, it is also fruitful to analyse the expected value of $\#(KI_{at})$, given by:

$$\mathbb{E}[\#(KI_{at})] = \sum_t [(1 - I_{at})I_{ct}.P(I_{ct} = 1|I_{at} = 0) + (1 - I_{at})(1 - I_{ct}).P(I_{ct} = 0|I_{at} = 0)] \quad (8)$$

which, because $KI_{at} \subset KC_{ct}$, gives $P(I_{ct} = 1|I_{at} = 0) = P(I_{ct} = 1) - P(I_{at} = 1)$ and $P(I_{ct} = 0|I_{at} = 0) = 1 - P(I_{ct} = 1)$, which simplifies $\mathbb{E}[\#(KI_{at})]$ to

$$\mathbb{E}[\#(KI_{at})] = \sum_t [1 - P(I_{at} = 1)] \quad (9)$$

i.e. individual a learns something new according to the probability that she is not satisfied with her current knowledge when prompted with phenomenon $phen_{at}$.

If we follow the same path, we can also define a cardinality and an expected cardinality function for KC_{ct} , which, given the algorithmical solution to the individual's problem (i.e. an agent first looks at KI_{at} , then at KC_{ct} , then uses $create(phen_{at})$ and given $KI_{at} \in KC_{ct}$, must yield:⁹

$$\mathbb{E}[\#(KC_{ct})] = \sum_t [1 - P(I_{ct} = 1)],$$

i.e. group c learns something new according to the probability that individual a is not satisfied with his group's current knowledge when prompted with phenomenon $phen_{at}$.

⁹Given that the possibility space is equal to $\Omega = I_{at}I_{ct} + (1 - I_{at})I_{ct} + I_{at}(1 - I_{ct}) + (1 - I_{at})(1 - I_{ct})$, that $P(\Omega) = 1$, and that, because $KI_{at} \in KC_{ct}$, $P[I_{at}(1 - I_{ct})] = 0$, we get that $P[(1 - I_{at})(1 - I_{ct})] = 1 - (P[I_{at}I_{ct}] + P[(1 - I_{at})I_{ct}]) = 1 - P[I_{ct}]$.

Moreover, if we look at the comparative statics for $\mathbb{E}[\#(KI_{at})]$ and $\mathbb{E}[\#(KC_{ct})]$, we see that $\frac{\partial \mathbb{E}[\#(KC_{ct})]}{\partial t} \geq 0$, $\frac{\partial^2 \mathbb{E}[\#(KC_{ct})]}{\partial t^2} < 0$, $\frac{\partial \mathbb{E}[\#(KC_{ct})]}{\partial s} \leq 0$, $\frac{\partial^2 \mathbb{E}[\#(KC_{ct})]}{\partial t \partial s} \leq 0$ and $\frac{\partial \mathbb{E}[\#(KI_{at})]}{\partial t} \geq 0$, $\frac{\partial^2 \mathbb{E}[\#(KI_{at})]}{\partial t^2} < 0$, $\frac{\partial \mathbb{E}[\#(KI_{at})]}{\partial s} \leq 0$, $\frac{\partial^2 \mathbb{E}[\#(KI_{at})]}{\partial t \partial s} \leq 0$.

In the following graph (figure 1), we plot the cardinality of the collective knowledge stock through times, in a simulation of the model for varying satisfaction parameters, maintaining the noise radius fixed. More specifically we varied satisfaction at incremental rates of 0.01 from 2.0 to 3.0, from 4.5 to 5.5, from 7.0 to 8.0 from 9.5 to 10.5, for 1000 iterations each, while maintaining the noise radius (i.e. the precision parameter) fixed at 5.0; for 20 agents, which were all part of the same collective.

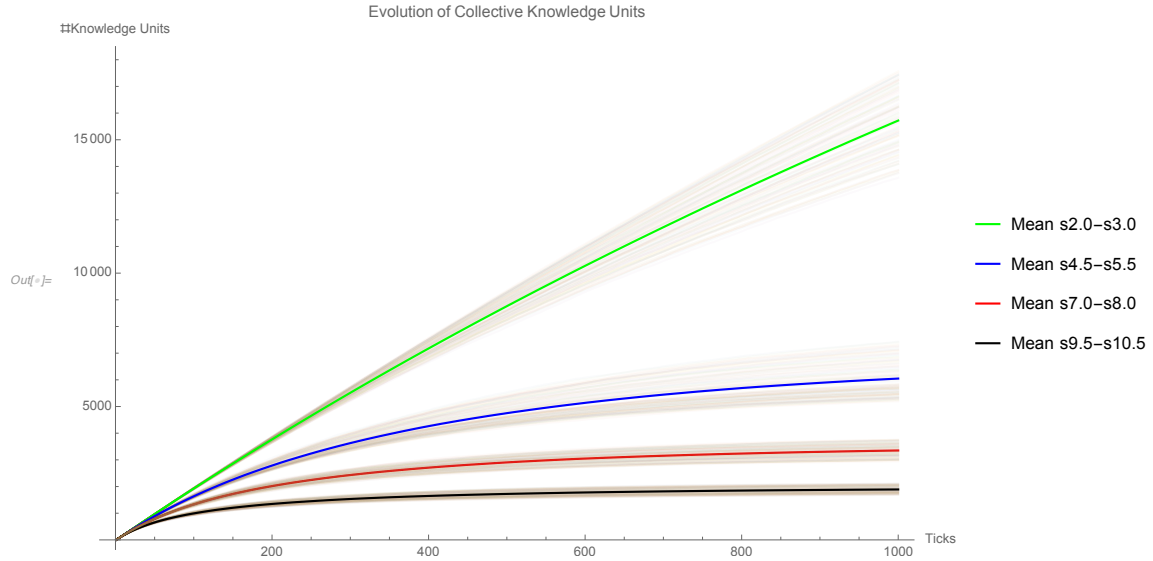


Figure 1: Cardinality of collective knowledge in a single collective of 20 agents, with 5000 iterations.

In the bottom line, with satisfaction parameter set around 10, there is the least amount of total knowledge units in the collective knowledge stock, followed by the values set around 7.5, then 5, then 2.5. It is possible to see that in the middle of the progression from 4.5 to 5.0 there is a change in the degree of variability. This seems to happen around when the satisfaction parameter surpasses the precision parameter at 5, which would be in line with what is expected from the model.

Below, in figure 2, we increased the number of iterations to 5000, and the same behavior can be observed. A further inquiry into of what is occurring will be provided in the next session, when we discuss the effects of the dynamic selection of

phenomena.

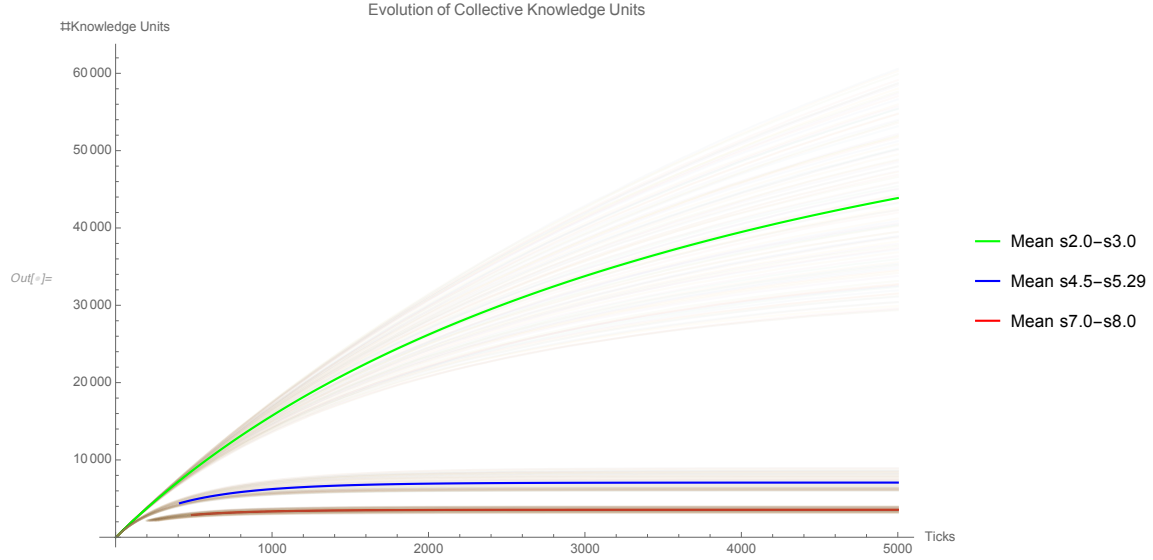


Figure 2: Cardinality of collective knowledge in a single collective of 20 agents.

The overall observable characteristic is that for an increasing s , i.e. for a relaxation in how demanding we are of the knowledge that we individually possess or that we may acquire from the collective knowledge stock, the amount of knowledge produced decreases. Which, again, is in line with what the comparative statics of the expected cardinality of KC_{ct} predict.

4.2 The initial distribution of phenomena in space

If we consider the distribution of euclidean distances between all phenomena, the higher the occurrence of distances smaller than $r + s$, the higher the probability that a knowledge point will be able to explain multiple phenomena. For instance, any $phen_{jt'}$ ($\forall jt' \neq it$) located at a distance greater than $r + s$ apart from $phen_{it}$ has probability zero of being satisfactorily explained by k_{it} . That is, the initial distribution of phenomena in the reality space is fundamental to the analysis of the model.

However, even more important (and interesting) than the initial distribution of distances between all the phenomena is the dynamic series of phenomena that occurred to each agent and to each collective. Interesting because it is only knowledge formulated to explain previous phenomena that can be used to explain any present

phenomenon.¹⁰ For instance, even if there are many phenomena clustered together at a distance smaller than $r + s$, the first of these to be selected will have a much smaller probability of being explained by existing knowledge than the following ones.

To illustrate this dynamic aspect, the following graphs (in figure 3) plot the evolution of the distance between a selected phenomenon and the nearest from all previous phenomena (on the top row). It is possible to see that the more previous phenomena there are (i.e. higher tick count), the smaller these distances to the nearest are. The red horizontal bar in each of these graphs represent the sum $r + s = \{7.5, 10.5, 12.5\}$, from left to right. I plot these bars to highlight the mechanism explored in the previous paragraphs, where the value of $r + s$ (i.e. noise + satisfaction) influences the probability that a phenomenon will be explained by existing knowledge.

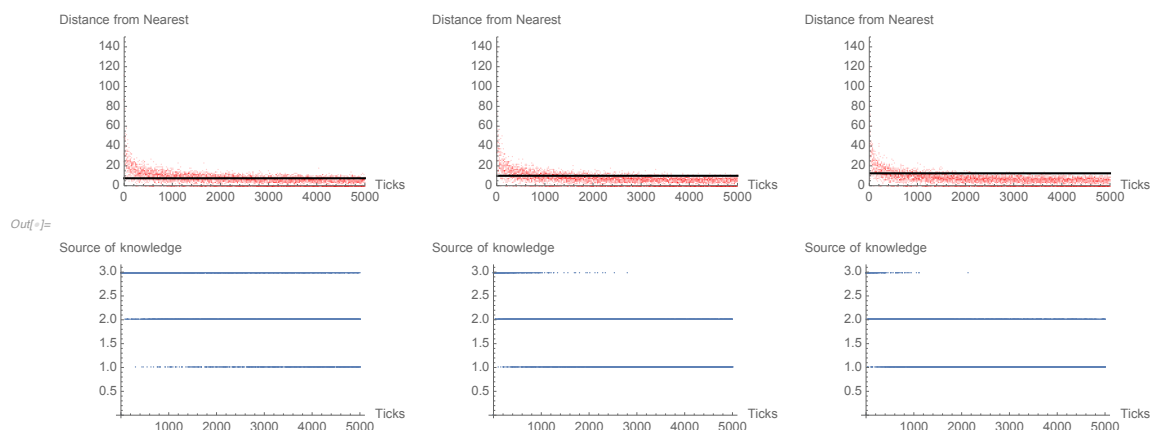


Figure 3: Time evolution of distances from the nearest (top row) and time evolution of sources of knowledge (bottom row) - for $p = 5$ and $s = \{2.5, 5.0, 7.5\}$, from left to right.

The graphs in the bottom row complete the argument showing the source of knowledge used to deal with any one phenomenon at any given time/tick t - with 3 representing creation, 2 representing collective knowledge stock and 1 representing individual knowledge stock.

The way to think about the figure is in vertical pairs, which are different from each other only in respect to the value of the satisfaction parameter s , which changes from 2.5, 5.0 and 7.5, from left to right, making the $r + s$ sums equal to 7.5, 10.0 and

¹⁰With this assumption I do not believe we rule out the possibility that knowledge used to explain previous imaginary phenomena, such as those portrayed in science fiction, can be used to explain a present phenomenon, imaginary or not.

12.5.

Following, in figure 4, I plot the histograms for each of the pairs of figures, where we can see that for each increase in the value of s , the sources of knowledge change from more creation to more individual, even though the distribution of distances from the nearest remain extremely similar.

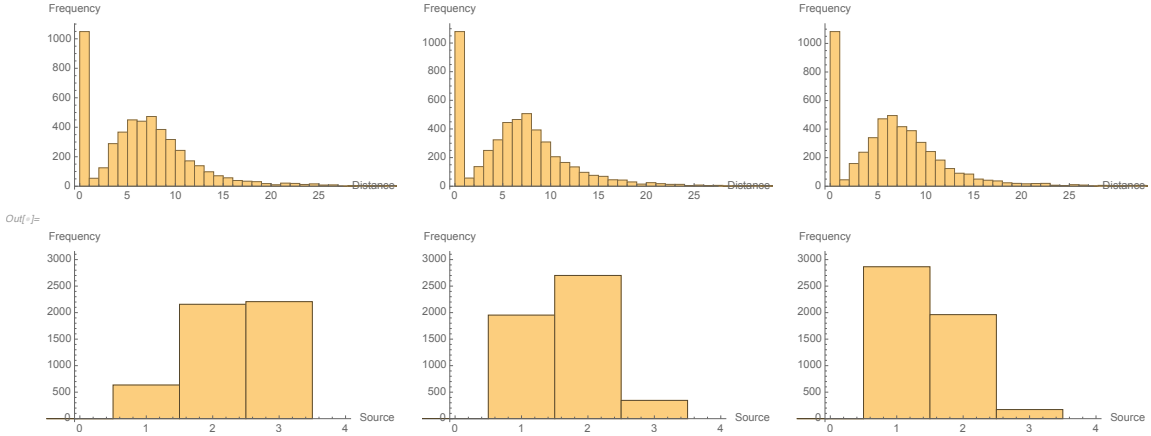


Figure 4: Histogram of distances from the nearest (top row) and histogram of sources of knowledge (bottom row) - for $p = 5$ and $s = \{2.5, 5.0, 7.5\}$, from left to right.

Since the only difference between the pairs is in the parameter s , we can argue that it is *causing* a change in the choice of knowledge source, which is a straight forward consequence of the model. What is less straight forward, however, is to find that the relation between the parameter s and the distances between the phenomena have an important role to play in the model.

It is obvious, but what the model is implying is that we can explain phenomena with the same understanding if they are closer together. We can use Newtonian gravity to explain the behavior of projectiles that obey certain rules - i.e. given that each is a manifestation of nature, a phenomenon, we can explain them with the same knowledge (Newtonian gravity), if they are closer together.

The evolution of knowledge, thus, is grounded in finding new ways in which phenomena can be clustered together, abstracting from their infinite set of characteristics, a subset of features that makes them comparable.

5 Conclusion

This paper proposes an abstract mechanism for the diffusion of knowledge and, with the use of computational methods, explores the logical ramifications of the hypothesis underlying the mechanism. In the model, all knowledge units bear a relation with all possible phenomena. This relation, intuitively named as *distance*, guides the usage of knowledge by each agent, in the following manner: if a knowledge unit is *too distant* it is not used, if it is the closest within a certain distance (s), it is used.

The agents follow a rather simple algorithmic procedure, where in each turn: agents interact with individual and collective knowledge stocks, and, if finding them insufficient to explain a phenomenon that is at hand, create new knowledge. As such, I propose a mechanism that provides a clear stopping rule for the creation of knowledge, as well as a rule for its creation.

Through simulations, I show that the aggregate quantity of knowledge units grows at a smaller rate, the bigger the parameter s . I also show that the similarity between phenomena plays an important role in the use of knowledge units.

In this specification, the distance s is the same for all agents in the economy. Future developments could expand the current methodology to account for idiosyncratic distances s_i , for each agent i . Moreover, the parameter s is exogenously defined, where it stands to reason that it should be a function of existing knowledge stocks. Lastly, there are multiple dimensions of knowledge, so it would be optimal if s were a multidimensional vector of acceptance parameters.

References

- R. Boschma. Proximity and innovation: a critical assessment. *Regional studies*, 39 (1):61–74, 2005.
- L. Casini and G. Manzo. Agent-based models and causality: a methodological appraisal, 2016.
- P. Gärdenfors. *Conceptual spaces: The geometry of thought*. MIT press, 2004.
- W. Harper. Newton’s methodology and mercury’s perihelion before and after einstein. *Philosophy of Science*, 74(5):932–942, 2007.
- P. Hedström and L. Udehn. Analytical sociology and theories of the middle range. *The Oxford handbook of analytical sociology*, pages 25–47, 2009.

- A. P. Kirman. Whom or what does the representative individual represent? *Journal of Economic Perspectives*, 6(2):117–136, June 1992. doi: 10.1257/jep.6.2.117. URL <http://www.aeaweb.org/articles?id=10.1257/jep.6.2.117>.
- A. Lieto, A. Chella, and M. Frixione. Conceptual spaces for cognitive architectures: A lingua franca for different levels of representation. *Biologically inspired cognitive architectures*, 19:1–9, 2017.
- M. P. Schlaile, J. Zeman, and M. Mueller. It’s a match! simulating compatibility-based learning in a network of networks. *Journal of evolutionary economics*, 28(5):1111–1150, 2018.
- M. Schlosser. Agency. In E. N. Zalta, editor, *The Stanford Encyclopedia of Philosophy*. Metaphysics Research Lab, Stanford University, fall 2015 edition, 2015.
- F. Squazzoni. *Agent-based computational sociology*. John Wiley & Sons, 2012.
- U. Wilensky. Netlogo flocking model. *Center for Connected Learning and Computer-Based Modeling, Northwestern University, Evanston, IL*, 1998.

UCD CENTRE FOR ECONOMIC RESEARCH – RECENT WORKING PAPERS

- [WP20/20](#) David Madden: 'The Socioeconomic Gradient of Cognitive Test Scores: Evidence from Two Cohorts of Irish Children' June 2020
- [WP20/21](#) Deirdre Coy and Orla Doyle: 'Should Early Health Investments Work? Evidence from an RCT of a Home Visiting Programme' July 2020
- [WP20/22](#) Cormac Ó Gráda: 'On Plague and Ebola in a Time of COVID-19' August 2020
- [WP20/23](#) Morgan Kelly: 'Understanding Persistence' September 2020
- [WP20/24](#) Jeff Concannon and Benjamin Elsner: 'Immigration and Redistribution' September 2020
- [WP20/25](#) Diego Zambiasi: 'Drugs on the Web, Crime in the Streets - The Impact of Dark Web Marketplaces on Street Crime' September 2020
- [WP20/26](#) Paul Hayes, Séin Healy, Tensay Meles, Robert Mooney, Sanghamitra Mukherjee, Lisa Ryan and Lindsay Sharpe: 'Attitudes to Renewable Energy Technologies: Driving Change in Early Adopter Markets' September 2020
- [WP20/27](#) Sanghamitra Mukherjee: 'Boosting Renewable Energy Technology Uptake in Ireland: A Machine Learning Approach' September 2020
- [WP20/28](#) Ivan Petrov and Lisa Ryan: 'The Landlord-Tenant Problem and Energy Efficiency in the Residential Rental Market' October 2020
- [WP20/29](#) Neil Cummins and Cormac Ó Gráda: 'On the Structure of Wealth-holding in Pre-Famine Ireland' November 2020
- [WP20/30](#) Tensay Meles and Lisa Ryan: 'Adoption of Renewable Home Heating Systems: An Agent-Based Model of Heat Pump Systems in Ireland' December 2020
- [WP20/31](#) Bernardo Buarque, Ronald Davies, Ryan Hynes and Dieter Kogler: 'Hops, Skip & a Jump: The Regional Uniqueness of Beer Styles' December 2020
- [WP21/01](#) Kevin Devereux: 'Returns to Teamwork and Professional Networks: Evidence from Economic Research' January 2021
- [WP21/02](#) K Peren Arin, Kevin Devereux and Mieszko Mazur: 'Taxes and Firm Investment' January 2021
- [WP21/03](#) Judith M Delaney and Paul J Devereux: 'Gender and Educational Achievement: Stylized Facts and Causal Evidence' January 2021
- [WP21/04](#) Pierluigi Conzo, Laura K Taylor, Juan S Morales, Margaret Samahita and Andrea Gallice: 'Can ❤️s Change Minds? Social Media Endorsements and Policy Preferences' February 2021
- [WP21/05](#) Diane Pelly, Michael Daly, Liam Delaney and Orla Doyle: 'Worker Well-being Before and During the COVID-19 Restrictions: A Longitudinal Study in the UK' February 2021
- [WP21/06](#) Margaret Samahita and Leonhard K Lades: 'The Unintended Side Effects of Regulating Charities: Donors Penalise Administrative Burden Almost as Much as Overheads' February 2021
- [WP21/07](#) Ellen Ryan and Karl Whelan: 'A Model of QE, Reserve Demand and the Money Multiplier' February 2021
- [WP21/08](#) Cormac Ó Gráda and Kevin Hjortshøj O'Rourke: 'The Irish Economy During the Century After Partition' April 2021
- [WP21/09](#) Ronald B Davies, Dieter F Kogler and Ryan Hynes: 'Patent Boxes and the Success Rate of Applications' April 2021
- [WP21/10](#) Benjamin Elsner, Ingo E Isphording and Ulf Zölitz: 'Achievement Rank Affects Performance and Major Choices in College' April 2021
- [WP21/11](#) Vincent Hogan and Patrick Massey: 'Soccer Clubs and Diminishing Returns: The Case of Paris Saint-Germain' April 2021
- [WP21/12](#) Demid Getik, Marco Islam and Margaret Samahita: 'The Inelastic Demand for Affirmative Action' May 2021