



Universiteit
Leiden
The Netherlands

Empirical signatures of universality, hierarchy and clustering in culture

Babeanu, A.I.

Citation

Babeanu, A. I. (2018, October 24). *Empirical signatures of universality, hierarchy and clustering in culture*. *Casimir PhD Series*. Retrieved from <https://hdl.handle.net/1887/66479>

Version: Not Applicable (or Unknown)

License: [Licence agreement concerning inclusion of doctoral thesis in the Institutional Repository of the University of Leiden](#)

Downloaded from: <https://hdl.handle.net/1887/66479>

Note: To cite this publication please use the final published version (if applicable).

Cover Page



Universiteit Leiden



The handle <http://hdl.handle.net/1887/66479> holds various files of this Leiden University dissertation.

Author: Babeanu, A.I.

Title: Empirical signatures of universality, hierarchy and clustering in culture

Issue Date: 2018-10-24

Empirical signatures of universality, hierarchy and clustering in culture

Proefschrift

ter verkrijging van
de graad van Doctor aan de Universiteit Leiden,
op gezag van Rector Magnificus prof. mr. C.J.J.M. Stolker,
volgens besluit van het College voor Promoties
te verdedigen op woensdag 24 oktober 2018
klokke 16.15 uur

door

Alexandru-Ionuț Băbeanu

geboren te Constanța (Roemenië)
in 1989

Promotor: Prof. Dr. J.M. van Ruitenbeek
Co-promotor: Dr. D. Garlaschelli

Promotiecommissie: Prof. Dr. G.T. Barkema (Universiteit Utrecht)
Prof. Dr. A. Flache (Rijksuniversiteit Groningen)
Prof. Dr. M. Verweij (Jacobs University Bremen, Duitsland)
Prof. Dr. J. Aarts
Prof. Dr. E.R. Eliel

The cover shows, in red, an example of a binary, rooted tree (dendrogram), mathematical construct which was very important for the research presented in this thesis. As a metaphor, the background shows a highly processed picture of a tree in Leidse Hout, taken by the candidate.

Casimir PhD series Delft-Leiden 2018-28
ISBN 978-90-8593-358-8

To my parents, Mariana and Laurențiu

Contents

1	Signs of universality in the structure of culture	7
1.1	Introduction	8
1.2	The formal representation of culture	10
1.3	Long-term cultural diversity and short-term collective behavior . .	12
1.4	Results	17
1.5	Discussion	21
1.6	Conclusion	24
1.A	Empirical data formatting	25
1.B	Feature-feature correlations	29
	Bibliography	31
2	Evidence for mixed rationalities in preference formation	37
2.1	Introduction	38
2.2	Model description	41
2.3	Model fitting	45
2.4	Model Outcomes	51
2.5	Discussion	55
2.6	Summary and conclusions	59
2.A	Controlling the generation of prototypes	60
2.A.1	Integer partition probabilities	60
2.A.2	Integer partition generation	62
2.B	Analytic calculations of model average inter-vector distance	64
2.C	Fitting algorithm	67
2.C.1	First level fitting	69
2.C.2	Second level fitting	72
2.C.3	Used functions	76
2.C.4	Algorithm usage	81
	Bibliography	83
3	Ultrametricity increases the predictability of cultural dynamics	87
3.1	Introduction	88
3.2	Ultrametricity and cultural dynamics	92

3.3	Partition-specific quantities	94
3.4	Predictability of the final state	97
3.5	Conclusion	100
3.A	Ultrametric-generation method	101
3.B	Detailed results	104
3.C	Dendrogram geometry	108
	Bibliography	110
4	A random matrix perspective of cultural structure	113
4.1	Introduction	114
4.2	Eigenvalue distributions for empirical data and null models	117
4.3	Two interpretations of structural modes	125
4.3.1	The feature-feature correlations scenario	128
4.3.2	The group structure scenario	129
4.3.3	Mathematical comparison of the two scenarios	131
4.4	Discriminating between the two interpretations	135
4.5	Revisiting the empirical data	141
4.6	Discussion	143
4.7	Conclusion	146
4.A	The fully-connected Ising (FCI) model	147
4.B	The symmetric two-groups (S2G) model	149
4.C	The structure of the FCI and S2G models	151
	Bibliography	152
	Samenvatting	155
	Summary	159
	Rezumat	163
	List of publications	166
	Curriculum vitæ	169
	Acknowledgments	171

Chapter 1

Signs of universality in the structure of culture

Understanding the dynamics of opinions, preferences and of culture as a whole requires more use of empirical data than has been done so far. It is clear that an important role in driving this dynamics is played by social influence, which is the essential ingredient of many quantitative models. Such models require that all traits are fixed when specifying the “initial cultural state”. Typically, this initial state is randomly generated, from a uniform distribution over the set of possible combinations of traits. However, recent work has shown that the outcome of social influence dynamics strongly depends on the nature of the initial state. If the latter is sampled from empirical data instead of being generated in a uniformly random way, a higher level of cultural diversity is found after long-term dynamics, for the same level of propensity towards collective behavior in the short-term. Moreover, if the initial state is randomized by shuffling the empirical traits among people, the level of long-term cultural diversity is in-between those obtained for the empirical and uniformly random counterparts. The current study repeats the analysis for multiple empirical data sets, showing that the results are remarkably similar, although the matrix of correlations between cultural variables clearly differs across data sets. This points towards robust structural properties inherent in empirical cultural states, possibly due to universal laws governing the dynamics of culture in the real world. The results also suggest that this dynamics might be characterized by criticality and involve mechanisms beyond social influence.

This chapter is based on the following scientific article:
A. I. Băbeanu, L. Talman and D. Garlaschelli *Eur. Phys. J. B* 90: 237 (2017).

1.1 Introduction

Quantitative, interdisciplinary research on social systems has recently seen a dramatic increase [1, 2], which is largely motivated by large amounts of data becoming available as a consequence of online and mobile phone activity. Such data sets allow one to map out large social networks [3, 4], consisting of connections and interaction patterns between humans, as well as to keep track of how these networks evolve with time [5]. This stimulated a series of empirical and theoretical studies of the structure and dynamics of social networks [6, 7, 8, 9]. Less attention has been paid to another, complementary aspect of social systems, having to do with the presence and evolution of opinions and preferences: the structure and dynamics of “culture”. This aspect particularly suffers from a lack of empirical research [10], which is what this article aims at partly compensating for.

This study makes use of quantitative tools developed within an interdisciplinary “cultural dynamics” research paradigm, which mostly consists of theoretical, model-driven studies, with significant input from physics [11]. In addition to embracing the dynamical nature of culture, this paradigm also embraces its multidimensional nature, although similar research focusing on single-dimensional dynamics also exists, in which case it is referred to as “opinion dynamics” [11] – interesting parallels between opinion dynamics and statistical physics were pointed out already in Ref. [12]. For cultural dynamics, the so-called Axelrod model [13] is very representative. In this setting, an individual (or agent) is encoded as a sequence of cultural traits (opinions, preferences, beliefs) commonly referred to as a “cultural vector”. Every entry of the vector corresponds to one dimension of culture, also referred to as one “cultural variable” or one “cultural feature”. All vectors evolve in time, driven mainly by social influence interactions, along with other ingredients, depending on which version of the model is actually used [14, 15, 16, 17, 18, 19, 20, 21, 22]. Any such model requires that all traits of all agents in the initial state are somehow specified, which is usually done randomly, using a uniform probability distribution over the set of possible cultural vectors – a uniform “cultural space distribution”. This choice is natural if the aim is understanding the (effect of the) dynamics by means of the structure present in the final state, in the absence of any structure in the initial state.

Taking a somewhat different perspective, Refs. [23, 24] explored alternative classes of initial conditions, trying instead to understand the effect that the initial state has on the dynamics and on the final state. It became apparent that the final state is rather sensitive to the initial state. In particular, an initial state constructed from an empirical social survey behaved significantly different from an initial state that was generated in a uniformly random way [23]. This implies that cultural dynamics is sensitive to the structure inherent in empirical data. Such sensitivity is worth exploiting, in order to better understand the empirical structure. Thus, if the cultural vectors in the initial state correspond to real individuals, the outcome of social influence models can be used as a quantitative tool for gaining insight about how real individuals are distributed in cultural

space, and indirectly about cultural dynamics in the real world, since the initial cultural state can be regarded as a partial snapshot of the real world dynamics. This is, to a great extent, the perspective of the research presented here, which makes use of a quantitative technique developed in Ref. [23].

On one hand, this technique incorporates the idea of social-influence cultural dynamics, which is encoded by a measure of long-term cultural diversity (LTCD), which makes use of an Axelrod-type model [13] of cultural dynamics with a minimal set of ingredients. The LTCD quantity estimates the extent to which discrepancies between opinions survive after a long period of cultural dynamics governed by consensus-favoring social influence, in the absence of any other process. For any given set of cultural vectors (or cultural state), the values of LTCD are shown in correspondence with those of another quantity, which is a measure of short-term collective behavior (STCB). The STCB quantity estimates the propensity of the agent population to short-term coordination in terms of their opinions with respect to only one topic. This is done using a modification of the Cont-Bouchaud model [25] of social coordination, which employs, in a more implicit way, the idea of one-dimensional opinion dynamics driven by social influence, supposedly taking place on a much shorter time-scale. As described in Sec. 1.3, both the LTCD and the STCB quantities are, additionally, functions of the same free parameter, the bounded confidence threshold ω , which controls the maximal distance in cultural space for which social influence can operate. The common dependence on this parameter is what allows for LTCD to be plotted as a function of STCB.

On the other hand, this technique also incorporates the comparison between the empirical cultural state, a uniformly random cultural state and a shuffled one – the latter is constructed by randomly permuting the empirical traits among vectors, thus retaining only part of the empirical information. Each of the three cultural states induces, in the LTCD-STCB plot, a curve parametrised by the bounded confidence threshold. In Ref. [23], for the random cultural state, the curve was such that at least one of the two quantities attained a close-to-minimal value for any value of the bounded confidence threshold ω , meaning that STCB and LTCD were mutually exclusive. This apparently called for a more complicated description or otherwise suggested a paradox, since real-world societies seem to allow for both short-term collective behavior and long-term cultural diversity. However, for the empirical cultural state, the two aspects became clearly more compatible, with both quantities attaining intermediate values for a certain ω interval, which appeared a parsimonious way of reconciling LTCD and STCB. At the same time the shuffled state entailed a compatibility of LTCD and STCB which was intermediate between those obtained for the empirical and random states.

The current study is dedicated to checking the robustness of the LTCD-STCB behavior identified in Ref. [23] across different empirical data sets. As shown in Sec. 1.4, this behavior appears to be universal, robust across geographical regions and independent of the details of the feature-feature correlation matrix. These results are based on multiple sets of cultural vectors, constructed from several

empirical sources and examined using the technique briefly described above. The LTCD and STCB quantities employed by this technique are explained in more detail in Sec. 1.3. Moreover, Sec. 1.2 gives more details about the formalism behind “cultural states” and related concepts. Finally, Sec. 1.5 discusses the results presented throughout the study, possible criticism and questions that can be further investigated. The manuscript is concluded in Sec. 1.6. Note that, although the definitions in Sec. 1.2 and Sec. 1.3 are effectively the same as in Ref. [23], in view of their importance for this manuscript, they are explained again here from a somewhat different angle, while emphasizing certain aspects that previously were only implicit.

1.2 The formal representation of culture

The way a cultural state is encoded here is inspired by models of cultural dynamics, in particular by Axelrod-type models [13]. In this paradigm, one deals with a set of variables, called “cultural features”, which encode information about various properties that individuals can have, properties that are inherently subjective and that can change under the action of “social influence” arising during person-to-person interactions. By construction, these variables are allowed to attain only specific values which are here called “cultural traits”. The interpretation here is that cultural traits encode “preferences”, “opinions”, “values” and “beliefs” that people can have on various topics, where each topic is associated to one feature.

A “cultural space” consists of the set of all possible combinations of cultural traits entailed by the set of chosen cultural features, together with a measure of dissimilarity between any two combinations. Moreover, this dissimilarity, also called the “cultural distance”, is defined in such a way that it satisfies all the properties of a metric distance (non-negativity, identity of indiscernibles, symmetry and triangle inequality). The so-called “Hamming” distance is commonly employed for this purpose, which is meaningful as long as there is no obvious ordering of the traits of any feature. A cultural space is thus an abstract, discrete, metric space, where each point corresponds to a specific combination of traits. However, the cultural space is mathematically not a vector space, since there is no notion of additivity attached to it.

A cultural state is essentially the selection of points in the cultural space that needs to be specified for the initial state of cultural dynamics models. Such a selection is also referred to here as a “set of cultural vectors” (SCV), where one “cultural vector” is one possible combination of traits. Formally, this is not a set in the rigorous sense, but a multiset, since it may contain duplicate elements – identical sequences of traits. However, duplicate elements will rarely occur in the initial states constructed for this study, since the number of cultural vectors is in practice much smaller than the number of possible points of the cultural space. On the other hand, they will often occur in the final state. This manuscript uses “SCV” interchangeably with “cultural state”.

It is also convenient to consider the notion of “cultural space distribution” (CSD), as a discrete probability mass function taking the cultural space as its support. If the SCV is constructed in a uniformly random way, one implicitly assumes that the underlying cultural space distribution is constant – all combinations of traits are equally likely. If, however, the SCV is constructed from empirical data, the inherent structure may be thought to correspond to non-homogeneities in an underlying CSD, for which the data is representative.

Here, empirical SCVs are mainly constructed from social survey data. Cultural features are obtained from the questions that are asked in the survey, while the traits of each feature correspond to the possible answers associated to the question. Thus, a cultural vector represents a sequence of answers that one individual has given to the list of questions in the survey. Importantly, a question is selected and encoded as a feature only if it is reasonably subjective, meaning that it does not ask about demographic or physical aspects concerning the individual (like place of residence, marital status, age), and that every allowed answer should be plausible at least from a certain perspective of looking at the question, or for people with a certain background or a certain way of thinking. Moreover, a question is disregarded if the survey is defined in such a way that its list of a-priori allowed answers depends on what answers are given to other questions. All features remaining after this filtering – see Sec. 1.A of the Appendix for more details – are assumed to contribute equally to the cultural distance, but the way they contribute depends on whether they are treated as nominal or as ordinal variables. Specifically, the cultural distance d_{ij} between two vectors i and j is computed according to:

$$d_{ij} = \frac{1}{F} \sum_{k=1}^F \left[f_{\text{nom}}^k (1 - \delta(x_i^k, x_j^k)) + (1 - f_{\text{nom}}^k) \frac{|x_i^k - x_j^k|}{q^k - 1} \right] = \frac{1}{F} \sum_{k=1}^F d_{ij}^k, \quad (1.1)$$

where F is the number of cultural features with k iterating over them, f_{nom}^k is a binary variable encoding the type of feature k (1 for nominal and 0 for ordinal), q^k is the range (number of traits) of feature k , $\delta(a, b)$ is a Kroneker delta function of traits a and b (of the same feature) and x_i^k is the trait of cultural vector i with respect to feature k . This definition reduces to the Hamming distance in case there are only nominal variables present. The second equality sign gives a formulation of the cultural distance as a sum over feature-level cultural distance contributions d_{ij}^k/F .

These feature-level contributions allow one to formulate, following Ref. [23], a notion of feature-feature covariance:

$$\sigma^{k,l} = \frac{\langle d_{ij}^k d_{ij}^l \rangle_{i,j \in 1,N}^{i < j} - \langle d_{ij}^k \rangle_{i,j \in 1,N}^{i < j} \langle d_{ij}^l \rangle_{i,j \in 1,N}^{i < j}}{F^2} \quad (1.2)$$

valid for any two features k and l , regardless of f_{nom}^k and f_{nom}^l . Note that the averaging is performed over all $N(N-1)/2$ distinct pairs (i, j) , $i \neq j$ of cultural

vectors, rather than over all N cultural vectors. The feature-feature covariances can be used to define the associated feature-feature (Pearson) correlations via:

$$\rho^{k,l} = \frac{\sigma^{k,l}}{\sqrt{\sigma^{k,k}\sigma^{l,l}}} \quad (1.3)$$

which measures the extent to which large/small distances in terms of feature k are associated to large/small distances in terms of feature l . One can definitely see the $F \times F$ correlation matrix ρ as a reflection of a CSD that is compatible with the data. In general, however, the correlation matrix will only retain part of the information encoded in the CSD, first because $\rho^{k,l}$ retains only part of the information in the 2-dimensional contingency table of features k and l , second because a CSD is essentially an F -dimensional contingency table, which might entail all kinds of higher-order correlations.

Assuming the definition of cultural distance given by Eq. (1.1), a cultural space is already specified by the list of features taken from an empirical data set, together with the associated ranges and types. In this empirically-defined cultural space, it is meaningful to talk about several types of SCVs. First, an empirical SCV is constructed from the empirical sequences of traits of the individuals selected from those sampled by the survey. Second, a shuffled SCV is constructed by randomly permuting the empirical traits among individuals, independently for every feature. Third, a random SCV is constructed by randomly choosing the trait of every person, for every feature. Note that the shuffled SCV exactly reproduces, for each feature, the empirical frequency of each trait, while disregarding all information about the frequencies of co-occurrence of various combinations of traits of two or more different features. Thus, shuffling destroys all feature-feature correlations $\rho^{k,l}$, as well as any higher-order correlations entailed by the empirical SCV, retaining only the information encoded in the marginal probability distributions associated to individual features. On the other hand, a random SCV retains nothing of the information inherent in the empirical SCV.

Finally, note that the mathematical definition of cultural distance illustrated by Eq. (1.1), already used in Refs. [23] in [24], is neither unique nor very sophisticated. Other definitions might capture differences in opinions, preferences, values, beliefs, attitudes and associated behavior tendencies in better, more precise ways – see Ref. [26] for a sophisticated approach. However, the current definition is arguably good enough for the problems explored in this study and for how they are attacked.

1.3 Long-term cultural diversity and short-term collective behavior

This section focuses on two quantities that are evaluated on sets of cultural vectors, namely the LTCD and STCB quantities mentioned above. These are based

on the ideas of cultural and opinion dynamics, respectively, driven by social influence in a population of interacting agents – as explained below, multidimensional cultural dynamics is explicitly implemented in LTCD, while unidimensional opinion dynamics is implicitly implemented in STCB. Each agent is associated to one of the cultural vectors in the SCV that is studied. For simplicity, both quantities assume that there is no physical space nor a social network that would constrain the interactions between agents. In both cases, the interactions are assumed to only be constrained by how the agents are distributed in cultural space. Specifically, only if the distance between two cultural vectors is smaller than the bounded confidence threshold ω are the two agents able to influence each other’s opinions in favor of local consensus: there needs to be enough similarity between the cultural traits of two people if any of them is to convince the other of anything. This picture is inspired by assimilation-contrast theory [27], Ref. [17] being the first study that explicitly uses the bounded-confidence threshold in the context of cultural dynamics, after having already been in use in the context of opinion dynamics for some time – see Ref. [28] for an overview. The bounded confidence threshold ω functions like a free parameter on which both the LTCD and the STCB quantities depend, for any given SCV.

The LTCD quantity is a measure of the extent to which the given SCV favors cultural diversity on the long term, namely a survival of differences in cultural traits at the macro level, in spite of repeated, consensus-favoring interactions at the micro level. In the real world, boundaries between populations belonging to different cultures appear to be resilient with respect to social interactions across them [29, 30, 31]. The measure relies on a Axelrod-type model [13] of cultural evolution with bounded confidence, which is applied on the SCV. This is meant to computationally simulate the evolution of cultural traits under the action of dyadic social influence, in the absence of other processes that may be present in reality. According to this model, at each moment in time, two agents i and j are randomly chosen for an interaction. If the distance d_{ij} between their cultural vectors is smaller than the threshold ω , then, with a probability proportional to $1 - d_{ij}$, for one of the features that distinguishes between the two vectors, one of the agents changes its trait to match the other. With time, agents become more similar to those that are within a distance ω in the cultural space. The dynamics stops when several groups are formed, within which agents are completely identical to each other, but too dissimilar across groups for any trait-changing interaction to occur. These groups are called “cultural domains”, term formulated in the context of the original Axelrod model [13], which also included a physical/geographical, 2-dimensional lattice but no (explicit) bounded confidence threshold. The normalized number of such cultural domains for a given value of ω , averaged over multiple runs of the model, defines the LTCD quantity:

$$\text{LTCD}(\omega) = \frac{\langle N_D \rangle_\omega}{N}, \quad (1.4)$$

where N_D is the cultural domains in the final (or absorbing) state of this model,

the normalization being made with respect to N , the size of the SCV.

The STCB quantity is a measure of the extent to which the given SCV favors collective behavior (or social coordination) on the short term, namely the extent to which the agents associated to the cultural vectors in the set would, due to social influence, tend to take actions or make choices in a similar, coordinated way rather than independently from each other. Bursts of fashion and popularity [32, 33, 34], rapid diffusion of rumors, gossips and habits [11, 35] and speculative bubbles and herding behavior on the stock markets [25, 36] are real-world examples of collective behavior on the short term. The measure relies on a Cont-Bouchaud type model [25], which deals with an aggregate choice or opinion of the entire agent population on one issue, which for simplicity is assumed here to be represented by a binary variable, which could encode, for instance, liking vs disliking an item. According to the model, when collectively confronted with this issue, the agents within a connected group effectively make the same choice or express the same opinion. In this context (where physical space and social network are disregarded), a connected group is a subset of agents that form a connected component in the graph obtained by introducing a link for every pair (i, j) of agents that are culturally close enough to socially influence each other $d_{ij} < \omega$. Based on this approximation, the aggregate, normalized choice of the entire population is expressed as a weighted average over the choices of the connected components, where the weight of the A th component is the size S_A of this component. However, the group choices themselves are still assumed to be binary, equiprobable random variables with values $\{-1, +1\}$. Thus, the aggregate, normalized choice is also a random variable, but one that is non-uniformly distributed over some set of rational numbers within $[-1, 1]$, in a manner that depends on the set of group sizes $\{S_A\}_\omega$ induced by a specific value of the ω threshold. The spread of this aggregate probability distribution provides the coordination measure that defines the STCB. It turns out that this quantity can be analytically computed, for a given ω , according to [23]:

$$\text{STCB}(\omega) = \sqrt{\sum_A \left(\frac{S_A}{N}\right)^2}, \quad (1.5)$$

where the summation is carried over the cultural connected components labeled by different A values. Note that only the sizes S_A of the components enter the calculation, which are in turn determined by the cultural graph obtained by thresholding the d_{ij} matrix by ω . Also note that STCB is higher when the agents are more concentrated in fewer and larger components.

There is a crucial difference between the LTCD and the STCB measures: while the former assumes that agents move in cultural space under the action of social influence, the latter assumes that the agents remain fixed in cultural space while they make their decision on one issue which is external to the cultural space. Although the STCB implicitly assumes that social influence occurs within the

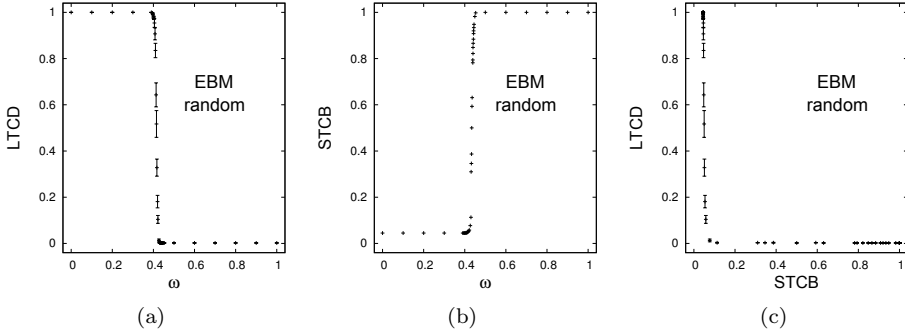


Figure 1.1: The interplay between long-term cultural diversity and of short-term collective behavior for a random set of cultural vectors. Showing the $LTCD(\omega)$ dependence (a), the $STCB(\omega)$ dependence (b) and the ω -induced $LTCD$ - $STCB$ correspondence (c), for a random set of $N = 500$ cultural vectors, in the cultural space of the Eurobarometer (EBM) data set (see Sec. 1.4).

cultural components, this influence is supposedly too superficial and too short-lived too also alter the cultural vectors themselves. Thus, the $LTCD$ and $STCB$ quantities are concerned with two different time-scales: a long time-scale for which cultural vectors and distances are dynamic and a short time-scale for which cultural vectors and distances are fixed. Moreover, while $LTCD$ requires computer simulations, the $STCB$ is computed in an analytical way. Thus, $LTCD$ can be seen as a characteristic of the final cultural state resulting from a long, cultural dynamics process, while the $STCB$ can be seen as a property of the initial cultural state.

It is worth explicitly illustrating, with Fig. 1.1, the behavior of the $LTCD$ and the $STCB$ quantities for a random SCV. The SCV is defined with respect to the cultural space of one of the data sets introduced in Sec. 1.4. Figs. 1.1(a) and 1.1(b) show, respectively, the dependence of the $LTCD$ and $STCB$ measures on the bounded-confidence threshold ω , while Fig. 1.1(c) shows the correspondence between the $LTCD$ and $STCB$ measures obtained by eliminating ω . The same data points are used for all 3 plots, where each point records all the 3 quantities ($LTCD$, $STCB$ and ω). The $LTCD$ quantity is averaged, for each point, over 10 runs of the cultural dynamics model, with the associated standard deviations shown by the error bars.

Fig. 1.1(a) shows that $LTCD$ decreases with ω : for large N , $LTCD$ goes from 1 to 0 as ω goes from 0 to 1. This is due to ω controlling the range of interaction in the cultural space. In general, convergence of agents happens in parallel in several regions of the cultural space, towards several points that are out of range of each other. Thus, ω also controls the expected number of such convergence points, which in turn determines the expected number of cultural domains in the

final state and thus the LTCD value – the latter three quantities decrease with increasing ω . If ω is small enough, there is effectively no successful interaction and thus no movement in cultural space, so each agent “converges” to one, distinct point (assuming that all vectors are different from each other in the initial state). If ω is large enough, all agents tend to converge to the same point in the cultural space. Note that, in terms of ω , these two extreme cases are actually two regimes, separated by a sharp decrease of LTCD over some intermediate ω interval. This sharp decrease can actually be understood as an order-disorder phase transition, where the disordered phase corresponds to low ω , while the ordered phase corresponds to high ω . This type of transition has been previously studied in the context of the Axelrod model [37, 21], although in terms of a differently defined control parameter – the (average) feature range q rather than the bounded-confidence threshold ω .

Fig. 1.1(b) shows that STCB is decreasing with ω : in the limit of large N , STCB goes from 0 to 1 as ω goes from 0 to 1. This is due to ω controlling the extent to which agents are culturally connected to each other. Higher ω implies fewer, but larger connected components in the cultural graph, thus a higher predisposition for coordination. If ω is small enough, there is one connected component for every agent, while if ω is small enough, there is one connected component containing all agents. Similarly to above, these two cases correspond to two regimes separated by a sharp increase of STCB, which can be again understood as a phase transition – it is actually a symmetry breaking phase transition, as explained in Ref. [23].

Fig. 1.1(c) shows that, as ω increases, one goes from the upper-left corner (high LTCD, low STCB) to the lower-right corner (low LTCD, high STCB), by first passing through the lower-left corner (low LTCD, low STCB). In other words, the sharp decrease of LTCD happens before the sharp increase of STCB, meaning that the critical ω of the LTCD phase transition is lower than that of the STCB phase transition. This is also visible at a close, comparative inspection of Figs. 1.1(a) and 1.1(b). The ω -region for which both the LTCD and the STCB attain low values corresponds to a special situation for which there is a relatively high level of convergence in the final cultural state (low LTCD), in spite of a relatively low level of connectivity in the initial cultural state (low STCB). This is apparently explained by the fact that movement in cultural space at a certain point in the cultural dynamics simulation facilitates further movement that would not have been possible at an earlier moment, so it is enough to have a few pairs of agents that can initially influence each other to gradually set a large fraction of the other agents in motion and in the end achieve a large amount of convergence. In any case, Fig. 1.1(c) shows that at least one of the two quantities has to attain a close-to-minimal value, regardless of the bounded-confidence threshold ω .

According to the considerations above, long-term cultural diversity and short-term collective behavior seem to be mutually exclusive, suggesting a paradox [23], at least if one accepts that real socio-cultural systems allow for both aspects. However, the above calculations make use of a random SCV, which assumes that the underlying cultural space distribution is uniform. Ref. [23] showed that an

empirical SCV allows for much more compatibility, with both quantities attaining intermediate values for a certain ω interval – as shown in Sec. 1.4, this translates to a higher LTCD-STCB curve than the one shown in Fig. 1.1(c) – meaning that the apparent paradox is solved by using realistic data about cultural traits. Moreover, a shuffled SCV entails a compatibility level that is in between those entailed by a random and by an empirical SCV. Thus, Ref [23] showed that an empirical SCV has enough structure to dramatically affect the behavior of social-influence dynamics acting upon it, aspect which had been neglected in the past.

1.4 Results

The findings of Ref. [23] are based on one data set. It is important to understand whether the observed properties are in fact robust across different populations and across different topics. This is accomplished by repeating the analysis of Ref. [23] on four data sets. These are taken from different sources, thus containing different cultural features and recording the traits of different people. The four data sources are: the Eurobarometer (EBM), containing opinions on science, technology and various European policy issues of people in EU countries [38]; the General Social Survey (GSS), containing opinions on a great variety of topics of people in the US [39]; the Religious Landscape (RL), containing religious beliefs and attitudes on certain political issues of people in the US [40]; Jester, containing online ratings of jokes [41].

Fig. 1.2 suggests that the properties highlighted by the LTCD-STCB curves are indeed universal. The 4 panels correspond to the 4 empirical data sets that are used. In each panel, the 3 curves correspond to the 3 levels of preserving the empirical information: full information (red), corresponding to the empirical SCV; partial information (blue), corresponding to the shuffled SCV; no information (black), corresponding to the random SCV. Note that, for every data set, the empirical SCV allows for more compatibility between LTCD and STCB than the shuffled SCV, which in turn allows for more compatibility than the random SCV. Also note that the empirical LTCD-STCB correspondence is always close to the second diagonal. These qualitative observations constitute the basis for the claim of there being universal structural properties underlying empirical sets of cultural vectors.

In relation to aspects discussed at the end of Sec. 1.2, the change of the LTCD-STCB curve when going from the random to the shuffled and further to the empirical CSV visible in Fig. 1.2 is related to the LTCD phase transition coming closer to the STCB phase transition. As ω increases, for the random case, the LTCD phase transition is almost over when the STCB phase transition begins, for the shuffled case there is more overlap between the high- ω part of the former and the low- ω part of the latter, while for the empirical case there is an almost perfect overlap between the two. The empirical behavior is illustrated by Fig. 1.3: within the $\omega \in [0.2, 0.4]$ interval, the decrease in LTCD is systematically accompanied

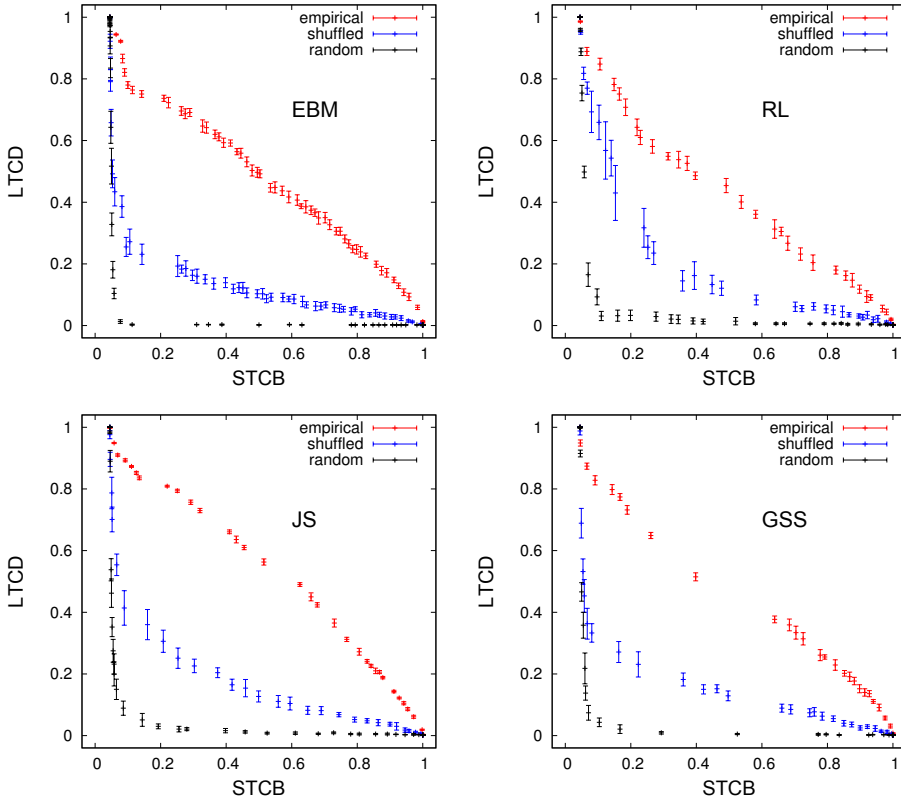


Figure 1.2: The correspondence between long-term cultural diversity (LTCD) and short-term collective behavior (STCB) for the empirical (red), shuffled (blue) and random (black) sets of cultural vectors, for four data sets: Eurobarometer (EBM), General Social Survey (GSS), Religious Landscape (RL) and Jester (JS). Error bars denote standard deviations over multiple cultural dynamics runs. There are $N = 500$ elements in each set of cultural vectors.

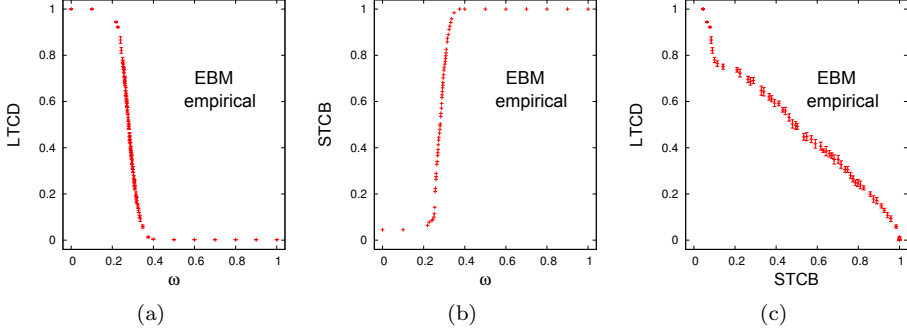


Figure 1.3: The interplay between long-term cultural diversity and short-term collective behavior for an empirical set of cultural vectors. Showing the LTCD(ω) dependence (a), the STCB(ω) dependence (b) and the ω -induced LTCD-STCB correspondence (c), for an empirical set of $N = 500$ cultural vectors, constructed from the Eurobarometer (EBM) data set.

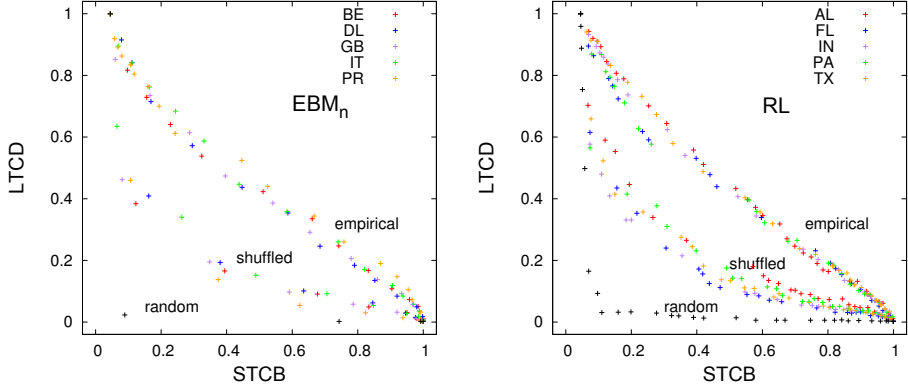


Figure 1.4: The correspondence between long-term cultural diversity (LTCD) and short-term collective behavior (STCB) for empirical and shuffled sets of cultural vectors constructed from country-level and state-level samples of Eurobarometer-nominal (EBM_n) data (left) and Religious Landscape (RL) data (right) respectively. There are $N = 500$ elements in each set of cultural vectors. For visual clarity, error bars are omitted and the same colors are used for both the empirical and shuffled cases, while the LTCD-STCB curve is also shown for one random set of cultural vectors in each case.

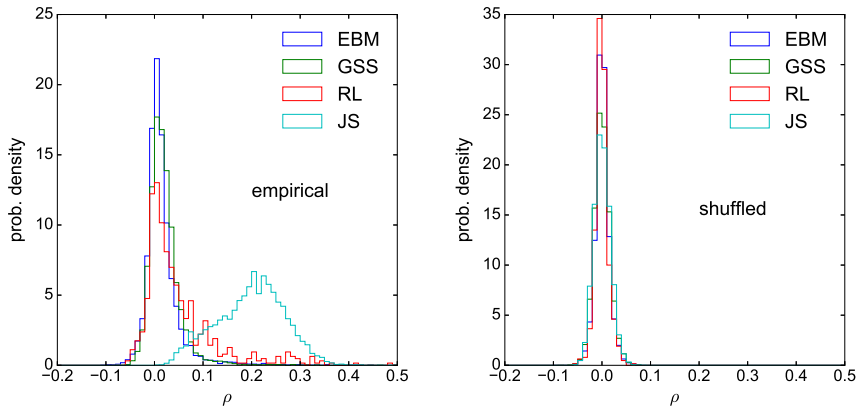


Figure 1.5: Distribution of feature-feature correlation ρ for the empirical (left) and shuffled (right) versions of each of the four data sets (legend). Each histogram is normalized such that its integral is equal to 1, after being initially filled with $F(F-1)/2$ entries, where F is the number of features in the respective data set, each entry corresponding to one pair (k, l) of distinct features. For the normalization, the integral multiplies the bin content with the bin width $\delta\rho$ (the same for all histograms): the ordinate value of each bin is its relative frequency multiplied by a factor of $1/\delta\rho$.

by an increase in STCB. If one accepts that real-world systems are favorable for both LTCD and STCB and that the respective quantities used here are defined in a sensible way, this reasoning suggests that real-world systems function close to criticality, from the perspective of both measures: only at criticality or close to it are both quantities allowed to attain non-vanishing values in the empirical case. In order to stay away from criticality, the system would need to abandon either the propensity towards LTCD or the propensity to STCB. This suggests, as a speculation or conjecture, that the concept of self-organized criticality [42] might actually play an important role in a complete theory of cultural dynamics. If this is correct, then a complete theory of cultural dynamics should have no need of fine-tuning the ω parameter.

Another important aspect is the robustness of the LTCD-STCB curves of Fig. 1.2 when switching from one geographical region to another, which is illustrated here by Fig. 1.4. This is done by focusing on the two data sets which allow for division of the sample in terms of geographical regions, namely the Eurobarometer and the Religious Landscape. Moreover, only the nominal-variable information in the Eurobarometer is being used, for reducing the computational time required to run the cultural dynamics model, as well as for illustrating the robustness of the results with respect to the sample of cultural variables that are

used. The empirical and shuffled LTCD-STCB curves are being shown for 5 EU countries (left) and for 5 US states respectively (right). Only one random curve is shown, because, for a specific data set, the country/state-level SCVs are defined with respect to the same cultural space, which is fully determined by the types and ranges of variables in the empirical data, which are the same regardless of the sample of people. Note that, for both data sets, the empirical and shuffled curves fall into clearly distinguishable bands. The empirical curves are systematically above the shuffled ones, while being again close to the second diagonal. This also suggests a geographical universality of the structural properties inherent in empirical data.

When confronted with these results, one thinks of unavoidable similarities between questions in the survey, which induce correlations between cultural features. Since these correlations are destroyed by the shuffling procedure, it is tempting to invoke them as an explanation for the discrepancy between an empirical LTCD-STCB curve and its shuffled counterpart. However, there is no reason to believe that such similarities are equally present in different empirical data sets, or that they are similarly distributed among the pairs of questions in the data set, since different data sets rely on completely different sets of variables. In fact, the measured feature-feature correlations $\rho^{k,l}$, defined via Eq. (1.3) are quite different across the four data sets used here. This is illustrated by Fig. 1.5, which shows how the values of these correlations are distributed for the different empirical SCVs (left), while also showing, for comparison, the distributions for their shuffled counterparts (right), which, as expected, are strongly peaked around 0 (the empirical and shuffled correlation matrices are shown in Figs. 1.6 and 1.7 of Appendix Sec. 1.B). The departure of the empirical distribution from its shuffled counterpart is clearly different across data sets, whereas the departure of the empirical LTCD-STCB curve from its shuffled counterpart is very similar across data sets, as shown in Fig. 1.2. Moreover, feature-feature correlations are typically small, given that any $\rho_{k,l}$ can take values within the $[-1, 1]$ interval. These are indications that the properties captured by the LTCD-STCB plot are not (or not exclusively) due to feature-feature correlations, and that additional information destroyed by shuffling (including higher-order correlations) plays an important role. Such considerations enforce the idea that the observed properties are due to a more subtle, dynamical and universal mechanism.

1.5 Discussion

The findings above stem from analyzing conventional social survey data in an unconventional way. Specifically, data from different sources is converted to empirical cultural states obeying a unified format, which does not retain the meanings of the questions in the survey, nor the meanings of their associated answers, but just the frequency distribution of respondents in cultural space. The LTCD and STCB quantities that are applied on the formatted data are also independent of

the meanings of used variables and values, although highly sensitive to the distribution in cultural space. This “semantically-invariant” nature of the analysis (invariance with respect to any relabeling of the cultural space that preserve all distances) is what allows one to potentially uncover universal properties in the structure of culture.

The results of the analysis suggest that there is something universal about how real people are distributed in cultural space. Empirical cultural states seem to induce a correspondence between LTCD and STCB that is highly robust across data sets, while significantly and consistently different from those induced by shuffled and random cultural states. If empirical cultural states are regarded as partial snapshots of this dynamics, the supposedly universal behavior could be seen as a consequence of general laws governing the dynamics of culture in the real world. This rises the question of what these laws actually are: what is the mechanism giving rise to distributions in cultural space that are compatible with the above results. Answering this question might mean achieving a full understanding of cultural dynamics. If one thinks in terms of snapshots of culture, this is equivalent to finding a general theory of preference formation, which is a fundamental challenge for the social sciences [43], with important implications for properly understanding decision making and economic behavior [44, 45, 46]. It appears that an important role for such a theory should be played by social influence, as its role in the aggregation of individual opinions and the formation of collective opinions has been extensively studied [12, 47, 48]. However, most of these studies focus on one-dimensional systems, while the empirical signatures presented are extracted from data with high dimensionality.

From a theoretical perspective, bringing together multidimensional opinion spaces and the notion of social influence is achieved by Axelrod-like models of cultural dynamics. Initializing the Axelrod dynamics with a random cultural state and studying the outcome goes along with understanding the type of structure that social influence can dynamically give rise to, assuming a structureless initial state. If social influence alone is responsible for the structure observed in empirical data, one would expect that an empirical cultural state is an intermediate outcome of the Axelrod dynamics. Thus, applying this dynamics to an empirical state would lead to an absorbing states that are statistically compatible with those obtained by applying the same dynamics to random states. However, the analysis presented here, whose LTCD quantity incorporates full simulations of an Axelrod-like model, shows a clear and robust discrepancy between the random and the empirical states. This suggests that social influence is not enough for explaining the generic empirical structure highlighted by the analysis. Nonetheless, the Axelrod model used by the LTCD quantity is highly simplistic, disregarding geographical space, social networks, influence of media and other aspects that are present in the real world. Moreover, the empirical cultural vectors correspond to individuals that are typically not interacting with each other directly in the real world, while they do so in the Axelrod model. Checking whether such considerations are sufficient for explaining the systematic discrepancies between random,

shuffled and empirical cultural states is an interesting topic for further research. If these are not sufficient, more exotic model ingredients should be considered, such as cognitive processes [49] or logical constraints across cultural features [50].

Contrary to the reasoning above, one can argue that the difference between the empirical and the shuffled regime of the LTCD-STCB analysis may simply be due to the presence of feature-feature correlations, which in turn are supposedly due to “design details” of the social survey, having to do with certain questions being similar to each other. Consequently, there would be no need to think about dynamical mechanisms responsible for the empirical structure. However, the a-priori expectation is that design-induced correlations are relatively weak: collecting social survey data is expensive, so the survey should be designed such that it captures as much as possible of the relevant degrees of freedom, by minimizing the similarities among questions. Moreover, remaining similarities should be specific to each data set, whereas the LTCD-STCB analysis gives highly similar results for different data sets. To better illustrate this counterargument, feature-feature correlations were measured in Sec. 1.4 and explicitly shown to be specific to each social survey, which is compatible with the idea that they largely depend on “design details” – see Appendix Sec. 1.B for more remarks along these lines. In fact, feature-feature correlations can be seen as one of several manifestations of a non-uniform cultural space distribution, which is certainly also affected by a-priori, survey-dependent similarities between features, but arguably not in an essential way. It is also worth noting that one cannot say to what extent a correlation between two features is caused by an a-priori similarity between the two questions and to what extent it arises dynamically due to the combination of processes taking place in the real world. One can even argue that trying to disentangle the a-priori contribution is entirely meaningless, partly because the questions themselves are formulated by humans who interact with each other and with society.

Another aspect that this study pointed out is the strong dependence of social influence cultural dynamics and its final outcome on the initial cultural state. This dependence becomes manifest in the analysis presented in Sec. 1.4 as the systematic departure of the LTCD-STCB curve corresponding to empirical data from those corresponding to the shuffled and random counterparts, confirming and expanding the results of Refs. [23, 24]. The dependence on initial states is rarely studied in the literature on cultural/opinion dynamics. A notable exception is Ref. [51]: upon analysing the Metropolis dynamics of the Ising model using an analytic technique developed in the context of opinion dynamics, a regime is found that allows for several, qualitatively different equilibrium states to be reached, depending on the initial configuration. It is also worth noting that, for studying the Axelrod model, Ref. [37] is using a non-uniform distribution in cultural space for randomly generating its initial states. Still, it is a distribution that can be factorized as a product of Poisson, feature-level distributions, encoding no structure in addition to that entailed by the feature-level non-uniformities. Refs. [23, 24] also suggest that initial state dependence can be understood in terms of an ultrametric

appearance of real cultural data, observation which Ref. [24] exploits for developing static models of cultural states characterised by a hierarchical organization in cultural space. Although this line of reasoning has not been used here, it should be further explored by future work.

Defining a (probabilistic) model of cultural states would be equivalent to specifying a cultural space distribution, the model being more realistic when the empirical data is better representative of this distribution. Such future research is further motivated by the robust behavior identified by this study, and by the observation that the three types of cultural states appear to roughly fall into three equivalence classes, in terms of the shapes of the associated LTCD-STCB curves. The purpose would be to design a model that generates artificial SCVs falling under the empirical equivalence class. Once the model is in place and properly tuned, the analysis of SCVs can in principle be extended to regimes that are not empirically accessible, due to limitations on F and N . This should allow for more detailed, statistical physics work to be done in relation to the phase transitions described in Sec. 1.3 and Sec. 1.4, such as finite-size scaling analysis and measurement of critical exponents. One might also achieve a better understanding of the extent to which the notion of self-organized criticality is important, by analysing the distribution of cluster sizes in cultural space for interesting ω values. At this point, this is highly speculative, based on the apparent complementarity between the LTCD and STCB transitions for empirical data, as well as on accepting that real-world systems are favourable for both long-term cultural diversity and short-term collective behavior. One can object by arguing that the shape of the LTCD and STCB transitions are sensitive to the exact mix of ingredients going in evaluating the two quantities – for instance, one can imagine using a more sophisticated Axelrod-type mode for evaluating LTCD. However, in the manner used here, LTCD and STCB are defined in a very similar, minimalistic way: adding more ingredients, such as geographical space and social networks, should be done in parallel for both quantities. It is plausible that additional ingredients would alter the two transitions in the same way, such that the relationship between LTCD and STCB is preserved.

1.6 Conclusion

This study is an additional step towards understanding the dependence of social-influence cultural dynamics on the initial cultural state. At the same time, it provides insights about the structure inherent in empirical cultural data by means of its effect on cultural dynamics, evaluated by the LTCD quantity, conditional on its effect on shorter time-scale opinion dynamics, evaluated by the STCB quantity. It turns out that the LTCD-STCB combination, together with comparisons between empirical data and randomized counterparts, suggest the existence of universal properties characterising how real people are distributed in cultural space. These properties seem to be present in spite of the variabilities of

the feature-feature correlation matrix across data sets. Further work is needed to understand in more depth the nature and implications of these properties.

Appendices

1.A Empirical data formatting

This section explains various details concerning the formatting of empirical data. As previously mentioned, four data sets were employed, each of which was collected by different entities, for different purposes and in different formats. In order for the analysis and modeling conducted here to be carried out consistently, the important information had to be extracted from each data set and expressed in one, unified format. Essentially, this format dictates that each data set has to provide a certain number of ordinal features and a certain number of nominal features, where each feature has a certain number of possible traits (the range q of the feature), and that the traits of every individual in the data set are recorded with respect to all these features. This unified format can be effectively thought of as a table of traits, where the rows correspond to the features and the columns correspond to the individuals. There are various challenges involved when converting the data into this format. It is worth explaining first the challenges that are more generic, relevant for several data sets and second the challenges specific to each data set.

One of the difficulties consists in deciding, for each variable, whether it should be used as cultural feature or not. The following is a (not entirely exhaustive) list of types of variables which are worth mentioning in this regard:

- demographic variables, such as those encoding “age”, “place of residence” or “ethnicity” are discarded, as they do not record subjective human traits;
- certain variables, that were not seen as demographic variables by the survey authors, are also discarded if they recorded information about something that is too much in the respondent’s past, or about something that cannot be easily related to subjective preferences, opinion, values, beliefs or behavioral tendencies that can be conceivably altered via social influence in a reasonably easy way; often, the boundary between what is subjective and what is objective not clear; nonetheless, one can strive to make these decision consistently at the level of every data set, which is what was done here;
- there are questions that ask opinions with respect to something that is differently defined for different people in the survey, such as: “how satisfied you are about how the the economy of this country is going recently?” – if

there are people from different countries in the data set, or “how satisfied you are with your life?”; these questions are also discarded;

- questions asking the respondent to self-evaluate a certain, personal trait, such as “would you say about yourself that you are more conservatory or liberal on political affairs”, are retained, assuming that the respondent mostly self-evaluates, in a reasonably objective way, a personal (subjective) trait, rather than expressing a subjective opinion about the personal trait;
- certain variables containing relevant information are also discarded if, due to the survey format, they can only be answered when certain answers are given to other variables, or if the set of possible answers explicitly depends on answers given to other questions, regardless of whether these “other” variables themselves are selected or not; including such variables would introduce inconsistencies in the encoding of cultural vectors, the definition of cultural distance and the shuffling and randomization procedures.

The variables that are retained for further analysis need to be encoded either as nominal or ordinal cultural features. Deciding between the two encoding options was done here using the following criterion: if there are more than two possible answers that are not “neutral” (see next paragraph) and they can all be conceivably ordered along the real axis, then the variable is encoded as ordinal; if, instead, there are only two answers (typically “Yes” and “No”) in addition to the neutral ones, or if the non-neutral answers cannot be ordered along the real axis in a consistent way, then the variable is encoded as nominal.

Most variables retained from the data sets also allow for one or more “neutral” answers (often called “missing values” in social science research, although this term usually is somewhat more general). These are usually labeled as “Don’t know”, “Refused” or “Not Answered”. For further analysis, these neutral answers are merged (if more than one are present). If the variable is to be encoded as nominal, neutral answers are mapped to one, additional cultural trait, side-by-side with traits originating from non-neutral answers. If the variable is to be encoded as ordinal, they are mapped to the middle of the ordinal scale – if there is an even number of possible answers, for each person, the choice is randomly made between the two answers closest to the middle of the scale.

Note that some data sets (GSS and EBM below) formally allow for another type of answer, labeled as “IAP” or “INAP” (inapplicable), which is here regarded as separate from neutral answers (although in social science research they are often all placed under the “missing values” umbrella term). IAP values are recorded, for certain respondents, when answers to a specific question are not expected from those respondents, for reasons having to do with the design of the survey. This happens for question that are only asked conditionally on answers given before. However, as mentioned above, these conditional variables are anyway discarded. Similarly, IAP values are also recorded for questions that are only asked to a certain sub-sample of the people, although not being conditional on some other

question, in which case those questions are either removed or, if the sub-sample is large enough, the formatting is restricted to it. Finally, IAP values are also recorded for split-ballot or split-form variables (see GSS and EBM explanations below), in which case specific procedures are followed, which effectively discard all IAP answers before further analysis. Thus, regardless of how exactly they occur, one does not need to map IAP answers to any trait, as they are all filtered out as a consequence of other formatting rules. Note that for the RL data set, although IAP answers are not explicitly mentioned anywhere, this could have been the case, since there are questions that are conditionally asked on other questions – instead of IAP answers, system-missing values are present in the SPSS file, typically marked by the “.” dot character.

First, this study made use of the **Jester 2 (JS)** data set [41], which consists of online ratings of jokes collected between November 2006 and May 2009. There are around 1.7 million continuous ratings (on a scale from -10.00 to +10.00) of 150 jokes from 59,132 users. For most users however, of the 150 jokes, only 128 are provided as items to be rated, as the other 22 were eliminated at a certain point in time. For this study, each of the 128 items is converted into an ordinal feature with 7 traits (by splitting the $[-10, 10]$ interval into 7 bins of equal size, while assuming that everything falling within one bin constitutes the same answer). Moreover, only the 2916 users that had rated all items were retained for further analysis – although this introduces some bias in the sample, one can argue that it is desirable to focus on individuals that have rated everything, as this is an indication of commitment on the respondent’s side.

Second, the research used the **Religious Landscape (RL)** data set [40], which consists of opinions and attitudes on various religious topics, but also on various political and social issues. These data were collected in 2007 via telephone interviews from all states of USA – this study only used the data obtained from the continental part of the USA (without Hawaii and Alaska). There are multiple questions asking about the religious affiliation of respondents, which were all discarded. This is partly based on the assumption that religious affiliation is closer to a demographic variable than to a feature that can be easily altered via social influence, partly based on the very large number of answers and the nested, hierarchical nature of how they are organized. For this study, 36 cultural features were constructed (18 nominal and 18 are ordinal), for a number of 35558 respondents.

Third, the research used the **Eurobarometer 38.1 (EBM)** data set [38], which consists of opinions on science, technology, environment and various EU political issues (mainly related to the open market and the economy). The data were collected during November 1992, from 12 countries of the EU, via face-to-face interviews. In this survey, there are several blocks of “coupled” variables which are all discarded: within each block, there are explicit internal constraints on how answers can be given (such as answering “yes” to at most 3 questions out of 8 that are available), which do not allow for a consistent encoding as a set of nominal or ordinal features.

Another challenge when formatting the EBM data set is posed by the split-ballot procedure: the sample of people is split into 2 ballots, and certain questions are asked in slightly different versions (small differences in formulation, answers listed in different orders etc.) to the two ballots, while both versions are present in the SPSS file for all individuals – for every respondent, an IAP answer is recorded for the version that is not used for that respondent. The most meaningful approach is to merge the two versions and eliminate all IAP answers – if both versions are kept, strong structural artifacts arise in the matrix of cultural distances [24]. Most of the split ballot variables are encoded as ordinal and have the same range (same number of non-neutral answers) in both versions, such that a one-to-one correspondence can be made, similarly to Ref. [24]. Some of them are still ordinal but have different ranges in the two versions. In all these cases, there is a difference of only one trait among the two versions, such that one range is an even number while the other is odd. In this case, the odd version is kept for the merging, which guarantees the existence of a middle trait to which all neutral answers can be directly assigned. The non-neutral answers from the even version are mapped to the closest answers in the odd version, in terms of the distance from the lowest-value answer, assuming that the distance between the lowest-value and highest-value answers is the same in the two versions (consistent with the definition of cultural distance in Eq. (1.1)). There is one split ballot variable which is encoded as nominal, in which case the difference consists in a second question being asked for one of the ballots, which is simply discarded. After all the formatting, 144 cultural features are constructed from this data set (54 nominal and 90 ordinal), for a number of 13026 respondents.

Fourth, the study used the **General Social Survey (GSS)** data [39], collected during 1993 in the USA via face-to-face interviews. The overall scheme of how questions are asked to respondents is arguably more complicated than for the EBM data set. First, there is a split-form procedure involved, which is equivalent to what is called “split-ballot” in the case of EBM: the respondents are split into two groups, with certain questions being asked in two, slightly different versions. All these questions are ordinal and have the same ranges in the two forms; they are handled like in the case of EBM. Independently of the split-form procedure, there is another procedure called “split-ballot”, which is methodologically somewhat different: the sample of respondents is split in 3 ballots (A,B,C), while some questions are only asked to 2 of the 3 ballots (A and B, B and C or A and C). This is handled by discarding the questions asked to only 2 of the 3 ballots. Independently of the split-ballot and split-form procedures, there is a set of questions, also used within the International Social Survey Program (ISSP), which are not asked to a small fraction of respondents (49 out of 1608 respondents). This is handled by discarding the 49 people not exposed to the ISSP questions. All in all, 133 cultural features are constructed from the GSS data (8 nominal and 125 ordinal), for a number of 1559 respondents.

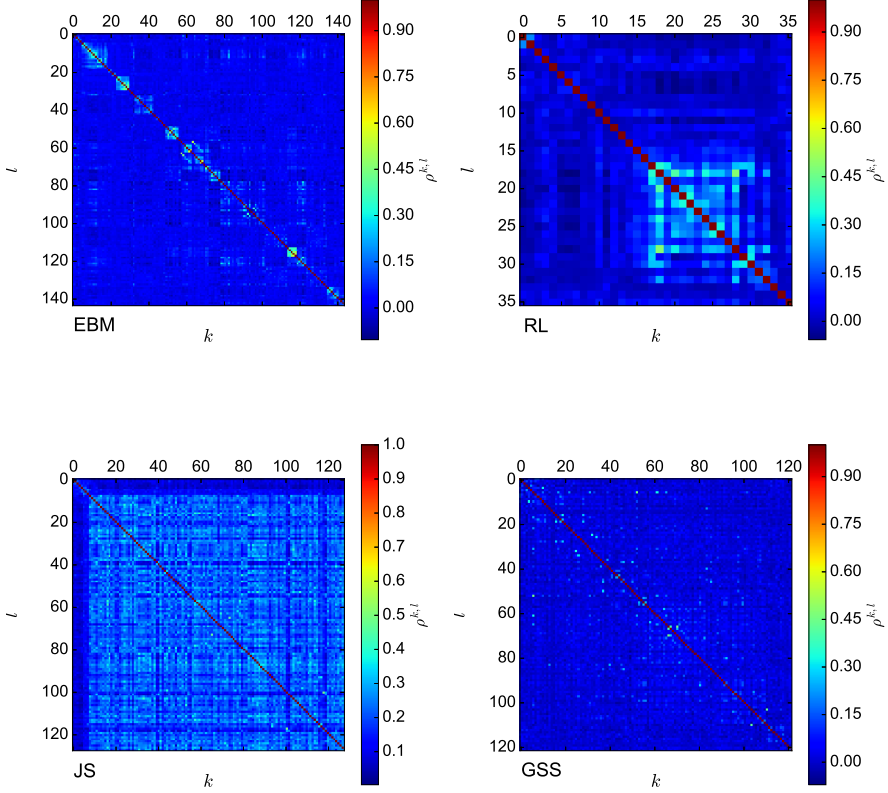


Figure 1.6: Matrix of feature-feature correlations in empirical sets of $N = 500$ cultural vectors obtained from the four sources: Eurobarometer (EBM), Religious Landscape (RL), Jester (JS) and General Social Survey (GSS). Each grid point shows the correlation $\rho^{k,l}$ between cultural features k and l .

1.B Feature-feature correlations

This section illustrates in detail the correlations between cultural features, computed according to Eq. (1.3). The feature-feature correlation matrices of the four empirical SCVs are shown in Fig. 1.6, while those of the four shuffled counterparts are shown in Fig. 1.7. The ordering of rows and columns is consistent with the actual ordering of questions in the four data sets. This leads to a partial block-diagonal aspect of the matrices associated to the Eurobarometer and Religious Landscape data sets, for which questions that deal with similar topics

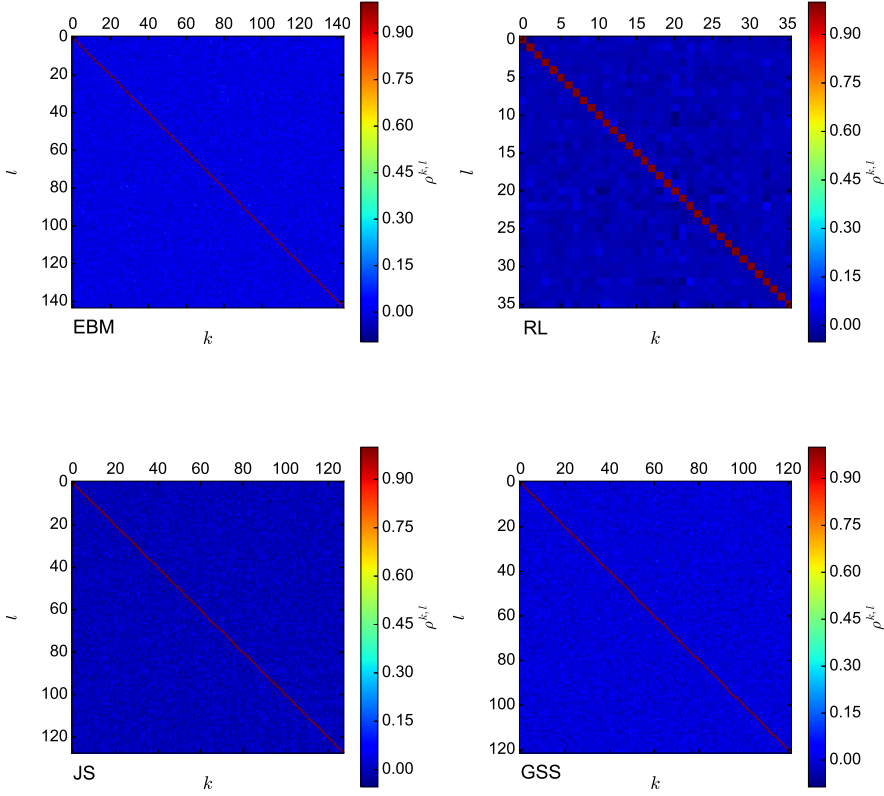


Figure 1.7: Matrix of feature-feature correlations in shuffled sets of $N = 500$ cultural vectors corresponding to the four empirical sources: Eurobarometer (EBM), Religious Landscape (RL), Jester (JS) and General Social Survey (GSS). Each grid point shows the correlation $\rho^{k,l}$ between cultural features k and l .

tend to appear next to each other. Note that, empirical correlations rarely show strong deviations from their shuffled counterparts. Interestingly, the largest level of correlation is visible for the Jester (JS) data set, which is certainly the least expensive to collect, since respondents provide their answers online, via an automated platform. Moreover, the second-largest level of correlation is present in the Religious Landscape (RL) data set, which is arguably the second-least expensive to collect, since it relies on telephone interviews, while the other two data sets rely on face-to-face interviews. This supports the idea that such correlations are survey specific, that they tend to be minimized by survey design and that they are not responsible for the generic structural properties identified by this study. There is a clear discrepancy between the Eurobarometer correlation matrix shown here and that shown in the Supplementary Information of Ref. [23]. However, the current study used a different, much more rigorous procedure of formatting the empirical data.

Bibliography

- [1] John Urry. The complexity turn. *Theory, Culture & Society*, 22(5):1–14, 2005.
- [2] David Lazer, Alex Pentland, Lada Adamic, Sinan Aral, Albert-László Barabási, Devon Brewer, Nicholas Christakis, Noshir Contractor, James Fowler, Myron Gutmann, Tony Jebara, Gary King, Michael Macy, Deb Roy, and Marshall Van Alstyne. Computational social science. *Science*, 323(5915):721–723, 2009.
- [3] Przemyslaw A. Grabowicz, José J. Ramasco, Esteban Moro, Josep M. Pujol, and Victor M. Eguiluz. Social features of online networks: The strength of intermediary ties in online social media. *PLOS ONE*, 7(1):1–9, 01 2012.
- [4] Jukka-Pekka Onnela, Jari Saramäki, Jörkki Hyvönen, Gábor Szabó, M Argollo de Menezes, Kimmo Kaski, Albert-László Barabási, and János Kertész. Analysis of a large-scale weighted network of one-to-one human communication. *New Journal of Physics*, 9(6):179, 2007.
- [5] Anuška Ferligoj, Luka Kronegger, Franc Mali, Tom A. B. Snijders, and Patrick Doreian. Scientific collaboration dynamics in a national scientific system. *Scientometrics*, 104(3):985–1012, Sep 2015.
- [6] Stanley Wasserman and Katherine Faust. *Social Network Analysis: Methods and Applications*. Cambridge University Press, 1994.
- [7] David Easley and Jon Kleinberg. *Networks, Crowds and Markets: Reasoning about a Highly Connected World*. Cambridge University Press, 2010.

- [8] Charles Kadushin. *Understanding Social Networks: Theories, Concepts and Findings*. Oxford University Press, 2012.
- [9] Petter Holme and Jari Saramäki. Temporal networks. *Physics Reports*, 519(3):97 – 125, 2012. Temporal Networks.
- [10] Pawel Sobkowicz. Modelling opinion formation with physics tools: Call for closer link with reality. *Journal of Artificial Societies and Social Simulation*, 12(1):11, 2009.
- [11] Claudio Castellano, Santo Fortunato, and Vittorio Loreto. Statistical physics of social dynamics. *Rev. Mod. Phys.*, 81:591–646, May 2009.
- [12] Serge Galam and Serge Moscovici. Towards a theory of collective phenomena: Consensus and attitude changes in groups. *European Journal of Social Psychology*, 21(1):49–74, 1991.
- [13] Robert Axelrod. The dissemination of culture. *Journal of Conflict Resolution*, 41(2):203–226, 1997.
- [14] Konstantin Klemm, Víctor M. Eguíluz, Raúl Toral, and Maxi San Miguel. Global culture: A noise-induced transition in finite systems. *Phys. Rev. E*, 67:045101, Apr 2003.
- [15] Konstantin Klemm, Víctor M. Eguíluz, Raúl Toral, and Maxi San Miguel. Nonequilibrium transitions in complex networks: A model of social interaction. *Phys. Rev. E*, 67:026120, Feb 2003.
- [16] Marcelo N. Kuperman. Cultural propagation on social networks. *Phys. Rev. E*, 73:046139, Apr 2006.
- [17] Andreas Flache and Michael W. Macy. Local convergence and global diversity: The robustness of cultural homophily. *arXiv:physics/0701333*, 2007.
- [18] J. C. González-Avella, M. G. Cosenza, and K. Tucci. Nonequilibrium transition induced by mass media in a model for social influence. *Phys. Rev. E*, 72:065102, Dec 2005.
- [19] Damon Centola, Juan Carlos González-Avella, Víctor M. Eguíluz, and Maxi San Miguel. Homophily, cultural drift, and the co-evolution of cultural groups. *Journal of Conflict Resolution*, 51(6):905–929, 2007.
- [20] Jens Pfau, Michael Kirley, and Yoshihisa Kashima. The co-evolution of cultures, social network communities, and agent locations in an extension of Axelrod’s model of cultural dissemination. *Physica A: Statistical Mechanics and its Applications*, 392(2):381–391, 2013.

- [21] Federico Battiston, Vincenzo Nicosia, Vito Latora, and Maxi San Miguel. Robust multiculturalism emerges from layered social influence. *CoRR*, abs/1606.05641, 2016.
- [22] Alex Stivala, Yoshihisa Kashima, and Michael Kirley. Culture and cooperation in a spatial public goods game. *Phys. Rev. E*, 94:032303, Sep 2016.
- [23] Luca Valori, Francesco Picciolo, Agnes Allansdottir, and Diego Garlaschelli. Reconciling long-term cultural diversity and short-term collective social behavior. *Proc Natl Acad Sci*, 109(4):1068–1073, 2012.
- [24] Alex Stivala, Garry Robins, Yoshihisa Kashima, and Michael Kirley. Ultrametric distribution of culture vectors in an extended Axelrod model of cultural dissemination. *Sci Rep*, 4(4870), 2014.
- [25] Rama Cont and Jean-Philippe Bouchaud. Herd behavior and aggregate fluctuations in financial markets. *Macroeconomic Dynamics*, 4(02):170–196, 2000.
- [26] Johannes Castner. Measures of cognitive distance and diversity. <http://ssrn.com/abstract=2477484>, 2014.
- [27] Muzafer Sherif and Carl Iver Hovland. *Social judgment: Assimilation and Contrast Effects in Communication and Attitude Change*. Yale University Press, New Haven, CT, 1961.
- [28] Jan Lorenz. Continuous opinion dynamics under bounded confidence: A survey. *International Journal of Modern Physics C*, 18(12):1819–1838, 2007.
- [29] Ruth García-Gavilanes, Yelena Mejova, and Daniele Quercia. Twitter ain’t without frontiers: Economic, social, and cultural boundaries in international communication. In *Proceedings of the 17th ACM Conference on Computer Supported Cooperative Work & Social Computing, CSCW ’14*, pages 1511–1522, New York, NY, USA, 2014. ACM.
- [30] Fredrik Barth. *Ethnic Groups and Boundaries*. Little, Brown and Company, Boston, 1969.
- [31] Robert Boyd and Peter J. Richardson. *The Origin and Evolution of Cultures*. Oxford University Press, New York, 2005.
- [32] Jukka-Pekka Onnela and Felix Reed-Tsochas. Spontaneous emergence of social influence in online systems. *Proc Natl Acad Sci*, 107(43):18375–18380, 2010.
- [33] Jacob Ratkiewicz, Santo Fortunato, Alessandro Flammini, Filippo Menczer, and Alessandro Vespignani. Characterizing and modeling the dynamics of online popularity. *Phys. Rev. Lett.*, 105:158701, Oct 2010.

- [34] Santo Fortunato and Claudio Castellano. Scaling and universality in proportional elections. *Phys. Rev. Lett.*, 99:138701, Sep 2007.
- [35] Bikas K. Chakrabarti, Anirban Chakrabarti, and Arnab Chatterjee. *Econophysics and Sociophysics: Trends and Perspectives*. Wiley-VCH Verlag GmbH & Co. KGaA, 2006.
- [36] Sitabhra Sinha, Arnab Chatterjee, Anirban Chakraborti, and Bikas K. Chakraborti. *Econophysics: An introduction*. Wiley-VCH Verlag GmbH & Co. KGaA, 2010.
- [37] Claudio Castellano, Matteo Marsili, and Alessandro Vespignani. Nonequilibrium phase transition in a model for social influence. *Phys. Rev. Lett.*, 85:3536–3539, Oct 2000.
- [38] Karlheinz Reif and Anna Melich. Euro-barometer 38.1: Consumer protection and perceptions of science and technology, november 1992. <https://doi.org/10.3886/ICPSR06045.v2>, 1995.
- [39] Tom W. Smith, Peter Marsden, Michael Hout, and Jibum Kim. General social surveys, 1993 ed. <http://gss.norc.org/get-the-data/spss>, 1972-2012.
- [40] Luis Lugo, Sandra Stencel, John Green, and Gregory Smith et al. U.s. religious landscape survey. religious beliefs and practices: Diverse and politically relevant. <http://www.pewforum.org/2008/06/01/>, 2008.
- [41] Ken Goldberg, Theresa Roeder, Dhruv Gupta, and Chris Perkins. Eigen-taste: A constant time collaborative filtering algorithm. *Information Retrieval*, 4(2):133–151, Jul 2001.
- [42] Per Bak, Chao Tang, and Kurt Wiesenfeld. Self-organized criticality: An explanation of the $1/f$ noise. *Phys. Rev. Lett.*, 59:381–384, Jul 1987.
- [43] Michael Thompson, Richard J. Ellis, and Aaron Wildavsky. *Cultural Theory*. Westview Press, 1990.
- [44] Ernst Fehr and Karla Hoff. Introduction: Tastes, castes and culture: the influence of society on preferences. *The Economic Journal*, 121(556):F396–F412, 2011.
- [45] Alain Cohn, Ernst Fehr, and M. A. Marechal. Business culture and dishonesty in the banking industry. *Nature*, 516(7529):86–89, 2014.
- [46] A. Cohn, J Engelmann, E. Fehr, and Michael André. Maréchal. Evidence for countercyclical risk aversion: An experiment with financial professionals. *American Economic Journal*, 105(2):860–85, 2015.

- [47] Serge Galam. Minority opinion spreading in random geometry. *The European Physical Journal B - Condensed Matter and Complex Systems*, 25(4):403–406, Feb 2002.
- [48] Serge Galam. Heterogeneous beliefs, segregation, and extremism in the making of public opinions. *Phys. Rev. E*, 71:046123, Apr 2005.
- [49] Pawel Sobkowicz. Opinion dynamics model based on cognitive biases. *arXiv:1703.01501v1*, 2017.
- [50] Noah E. Friedkin, Anton V. Proskurnikov, Roberto Tempo, and Sergey E. Parsegov. Network science on belief system dynamics under logic constraints. *Science*, 354(6310):321–326, 2016.
- [51] Serge Galam and André C. R. Martins. Two-dimensional ising transition through a technique from two-state opinion-dynamics models. *Phys. Rev. E*, 91:012108, Jan 2015.

Chapter 2

Evidence for mixed rationalities in preference formation

Understanding the mechanisms underlying the formation of cultural traits is an open challenge. This is intimately connected to cultural dynamics, which has been the focus of a variety of quantitative models. Recent studies have emphasized the importance of connecting those models to empirically accessible snapshots of cultural dynamics. In particular, it has been suggested that empirical cultural states, which differ systematically from randomized counterparts, exhibit properties that are universally present. Hence, a question about the mechanism responsible for the observed patterns naturally arises. This study proposes a stochastic structural model for generating cultural states that retain those robust empirical properties. One ingredient of the model assumes that every individual's set of traits is partly dictated by one of several universal "rationalities," informally postulated by several social science theories. The second, new ingredient assumes that, apart from a dominant rationality, each individual also has a certain exposure to the other rationalities. It is shown that both ingredients are required for reproducing the empirical regularities. This suggests that the effects of cultural dynamics in the real world can be described as an interplay of multiple, mixing rationalities, providing indirect evidence for the class of social science theories postulating such a mixing. The model should be seen as a static, effective description of culture, while a dynamical, more fundamental description is left for future research.

This chapter is based on the following scientific article:
A. I. Băbeanu and D. Garlaschelli, *Complexity*, Article ID 3615474 (2018).

2.1 Introduction

A solid theoretical understanding of how preferences form is currently lacking. There is little doubt that preferences, opinions, values and beliefs, which are generically referred to as “cultural traits”, are dynamical entities, and that interpersonal social influence plays an important role in driving their dynamics, among other factors. Moreover, a complete theoretical understanding should account for the fact that the dynamics of traits takes place in parallel along multiple dimensions, namely that opinions and preferences can develop in relation to multiple topics or aspects of life. Along these lines, various dynamical models been developed and studied [1], such as the Axelrod model [2], which is very representative for studies of multidimensional dynamics, commonly referred to as “cultural dynamics”, in contrast to studies of unidimensional dynamics, commonly referred to as “opinion dynamics”. Various studies of cultural dynamics extending the Axelrod model can be found in the literature [3, 4, 5, 6, 7, 8, 9, 10, 11]. Recent studies [12, 13, 14] (Chap. 1) have shown that models of cultural dynamics are sensitive to the initial conditions, namely to how the initial vectors of agents’ traits are chosen: initial cultural states constructed from empirical data show systematic deviations from their shuffled and random counterparts. In fact, Ref. [14] (Chap. 1) argues that such deviations point towards universal structural properties inherent in empirical cultural states. More insights about the formation of cultural traits should be achievable by studying these states, since they can be regarded as partial snapshots of cultural dynamics in the real world.

The universal properties mentioned above are expressed in terms of the effects the empirical cultural state has on social influence models whose initial conditions are specified by this state – here, a “cultural state” is a set of cultural vectors (SCV), where each cultural vector encodes the sequence of cultural traits associated to one agent in the model. On one hand, an Axelrod-type model [2] of (multi-dimensional) cultural dynamics is used to evaluate the propensity of the cultural state to long-term cultural diversity (LTCD). On the other hand, a Count-Bouchaud-type model [15] of (one-dimensional) opinion dynamics is used to evaluate the propensity of the cultural state to short-term collective behavior (STCB). Both measures are functions of a common parameter ω , controlling for the range of social influence in cultural space, which allows for an LTCD-STCB correspondence to be drawn for a given cultural state. It turns out that an empirical cultural state generally induces an LTCD-STCB curve that is close to the second diagonal ($\text{LTCD}(\omega) \approx 1 - \text{STCB}(\omega), \forall \omega$), while exhibiting, for a given STCB value, higher LTCD values than a trait-shuffled cultural state, which in turn exhibits higher LTCD values than a randomly generated counterpart [12, 14] (also see Chap 1). These results seem universal [14] (Chap. 1), namely independent of the data set used for constructing the cultural vectors composing the empirical cultural state, suggesting that real-world cultural dynamics is governed by universal laws. Moreover, as argued in Ref. [14] (Chap. 1), this type of analysis suggests that inter-agent social influence, the essential ingredient of cultural

dynamics models, is insufficient for explaining the observed structure. Although it is meaningful to incorporate additional ingredients into social influence models, while attempting to give rise to empirical-like structure in a dynamical setting, this study does not aim for that. Instead, it aims at providing an effective, phenomenological, static description of the observed structure, which should provide additional insights before developing a more fundamental, dynamical description.

The purpose of this study is to develop a structural stochastic model that would generate realistic cultural states, while incorporating plausible ingredients from social science. Specifically, these states should retain the universal properties inherent to empirical cultural states that are observed in Ref. [14] (Chap. 1). In fact, Ref. [13] has already investigated various ways of generating sets of cultural vectors in random, but non-uniform ways. A method that appeared particularly promising relied on the notion of “cultural prototypes”: a few underlying, abstract sequences of logically compatible, self-enforcing cultural traits, which govern the way the generated vectors are distributed in cultural space. According to the method, each cultural vector is partly a copy of one of the prototypes and partly random. The implicit claim is that each cultural prototype is induced by one of a few (3 to 5) fundamental and universal “principles of social life”, or “rationalities”, that would strongly affect any process of trait formation in any social system. Such entities are postulated, under different names and in slightly different numbers, by several theoretical frameworks in social science [16, 17, 18, 19, 20]. The exact number of such entities depends on the exact theory that is considered, as different theories are built on somewhat different arguments and pieces of evidence. It is important that the number is larger than 1 but not too large, while independent of system size. From a natural science perspective, such ideas are attractive, since they exhibit a certain reductionist tendency of trying to understand the observed socio-cultural variability in terms of combinations of a few, elementary and universal building blocks. Various parallels and similarities between these theories are discussed in the literature [21, 22, 23]. For the purpose of the current study, all these theories are equivalent. Still, for creating an instructive and compact context, the discussion is restricted to one of them, namely to Plural Rationality Theory, chosen for reasons discussed in Sec. 2.5.

Plural Rationality Theory (PRT), also referred to as “(Grid-Group) Cultural Theory” [16], is a qualitative description of socio-cultural structure and dynamics as an interplay between a small number of irreducible “ways of life”, or “rationalities”. These ways of life are understood as abstract, “elementary building blocks” of societies and are supposedly recognizable regardless of the geographical context, of the historical context or of the scale of the system that is studied. It is believed that the ways of life go along with different perceptions of risk [24, 25] and, interestingly, that they always coexist, although either of them is often dominant for a given period of time, for a given (part of) the system that one studies¹. Such

¹It may be useful to think of the ways of life as being the elements of a complete, orthogonal basis of some abstract vector space. One may then associate a vector in this space to a certain

ideas appear compatible with recent empirical findings concerning the existence of a small number of behavioural phenotypes in dyadic games [26]. In PRT, each way of life is understood as a self-enforcing combination of a “pattern of (social) relations” and a “cultural bias”. On one hand, a pattern of relations is often understood as a tendency of organizing the social ties between people in a certain way, thus a connectivity pattern in the social graph. On the other hand, a cultural bias is a combination of preferences, opinions, values and beliefs that are compatible with each other and with the associated pattern of relations. By comparison to the definitions in Ref. [14] (Chap. 1), one can easily realize that a cultural bias can be thought of as a point or a region in “cultural space” that is representative for the respective “way of life”. A cultural bias is formally represented here by the notion of “cultural prototype”, previously used in Ref. [13].

This notion is at the core of two stochastic, structural models of culture that are defined and studied here. The first model, called “Prototype Generation” (PG), postulates that each cultural vector is partly a copy of one of the k prototypes and partly random. This generation method is similar to the “Prototype Evolution” method of Ref. [13], although with small technical differences. The second model, called “Mixed Prototype Generation” (MPG), postulates that each cultural vector is an asymmetric mixture (or combination) of all the prototypes. From the perspective of PRT, this “mixing” is a formal realization of the idea that every person combines the ways of life in a unique way, such that preferences and opinions related to different aspects of life – cultural traits of different cultural features (or variables) – are due to the “influence” of different cultural biases, though at any given moment in time one cultural bias is usually dominating. In the literature concerned with PRT and the other, similar, theories, this mixing aspect often goes under the name of “the multiple self”, and was not implemented in Ref. [13]. The importance of mixing for correctly interpreting (and testing) PRT has been already stressed on [25], while the general importance of multiple selves for social science has also been extensively discussed [27]. Moreover, research on preferences in economic contexts also suggests that the multiple self is important [28, 29, 30]. On the other hand, research in cross-cultural psychology appears to be divided: some studies seem to ignore the multiple self [31], while others seem to acknowledge it [32, 33]. This study provides further insights on this matter, by directly comparing the PG and MPG models with each other and with empirical data,

Sec. 2.2 explains the models in detail, while Sec. 2.3 describes how the free parameters are tuned, as to reproduce some lower-order properties of one empirical cultural state. Cultural states generated with the two models are then evaluated, in Sec. 2.4, by means of the LTCD-STCB analysis of Refs. [12, 14] (Chap. 1). It is shown that cultural states generated by PG are structurally dissimilar to

part of a certain socio-cultural system, at a given moment in time. It is not clear to what extent such vectors would be related to the cultural vectors used in this study. This is only a semi-formal analogy that is not exploited further here, nor in any other study so far, to the extent that the authors are aware of.

the empirical ones, as they do not exhibit the universal LTCD-STCB behavior, after tuning the free parameters to empirical data in terms of simpler, but meaningful quantities. On the other hand, cultural states generated with MPG are structurally similar to the empirical ones, as they reproduce the universal LTCD-STCB behavior, after applying an analogous tuning procedure. This suggests that the mixing, multiple self ingredient is crucial for describing the effects of preference formation in terms of cultural prototypes, and that MPG should be regarded as the successful model. Sec. 2.5 further discusses the results, their limitations, as well as extensions of this work and questions that are worth investigating in the future. The manuscript is concluded in Sec. 2.6.

2.2 Model description

This section describes the two stochastic models of culture: the Prototype Generation (PG) model and the Mixed Prototype Generation (MPG) model, which are used below for generating sets of cultural vectors (SCVs) that can be quantitatively studied with the LTCD-STCB tool, previously applied to empirical SCVs in Refs. [12, 14] (Chap. 1). Both models rely on the concept of cultural prototype introduced above.

An SCV can be visualized as a table of cultural traits, where the columns correspond to cultural vectors (or sequences) and the rows correspond to cultural features (or variables). If the SCV is constructed from empirical data, the columns correspond to real people that are sampled by a social survey, while the rows correspond to questions that are asked in the social survey. This is illustrated by Fig. 2.1, which is explained in detail below. Consistently with Ref. [14] (Chap. 1), a “cultural space” is the set of all possible cultural vectors (or combinations of traits) allowed by the given set of cultural features: one combination of traits is one point in this discrete space. For the purpose of this work, the general set-up is restricted to cultural spaces defined in terms of features that are exclusively nominal. In this setting, distances between points in the cultural space are given by Eq. (2.5) of Sec. 2.3. Disregarding ordinal features makes the modeling paradigm compatible with the (arguably strong) assumption that one prototype corresponds to one point in cultural space, meaning that a prototype picks up one and only one trait of any given feature. Other limitations of this assumptions are extensively discussed in Sec. 2.5, together with possible ways of relaxing it, for the purpose of generalizing the current modeling paradigm in future work.

The two models are schematically illustrated in Fig. 2.1. The figure first shows a sketch of an empirical SCV, where the rows correspond to cultural features, the columns correspond to cultural vectors and the letters correspond to cultural traits – the n ’th row shows the traits of the N agents that are expressed (or formulated) with respect to the n ’th feature. Then, it shows a set of 3 cultural prototypes (their number could have been different), in 3 different colors, all of them spanning over all features (or questions) relevant for the empirical set of vectors. Finally,

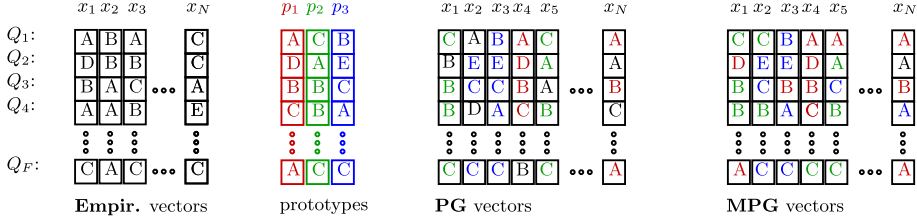


Figure 2.1: Schematic illustration of the two stochastic models, showing (from left to right): an empirical SCV with N vectors (x_1 to x_N) and F nominal variables (Q_1 to Q_F); a set of $k = 3$ cultural prototypes for the same F variables; a SCV with N vectors generated, from the prototypes, using the PG model; a SCV with N vectors generated, from the same prototypes, using the MPG model. For the PG and MPG sketches, red, green and blue denote the copies of cultural traits from one of the first, second and third prototype respectively, while black denotes the explicitly random generation of traits.

it illustrates a typical set of vectors generated using the PG method, followed by one generated using the MPG method. The colors distinguish between the prototypes, while indicating how the traits are copied from the prototypes to the cultural vectors, while black denotes traits that generated in an explicitly random way (uniform distribution, independently of the prototypes).

There are several things worth noting in relation to Fig. 2.1. First, the possibility that two or more prototypes pick the same trait for a certain feature is allowed by the current modeling paradigm (note that any of the traits that can be copied from one of the prototypes can also be generated via explicit randomness). This is essential for controlling the average prototype-prototype distance, as will become apparent below. Second, a PG vector is partly copied from one prototype and partly generated in an explicitly random way, while a MPG vector is a mixture of copies from all the prototypes, with one dominating prototype and with few traits generated in an explicitly random way. Third, both models make use of another type of randomness, in addition to the explicitly random trait generation and to the randomness involved in generating the prototypes. This randomness has to do with assigning every trait of every vector to a “prototype of origin”, once the random generation fraction and the influence fractions of the prototypes are specified. In the case of MPG, it is mainly this trait-assignment randomness that allows for the generation of a multitude of distinct cultural vectors from a small set of fixed prototypes, in the presence of little explicitly random trait generation.

The procedure for generating the cultural prototypes is the same for both the PG and the MPG models. One needs to specify the number of prototypes k , as well as the value of another parameter $\alpha \in (0, 1)$, which controls for the expected cultural distance between the prototypes. This parameter governs the expected

number of overlaps (or coincidences) between prototypes in terms of how they are distributed over the traits of a specific feature. In the extreme case of $\alpha \rightarrow 1$, all prototypes pick the same trait for every feature, yielding the smallest possible separation between the prototypes in cultural space (which coincides with the minimum of 0 allowed by the cultural distance definition in Eq. (2.5)). In the other extreme case of $\alpha \rightarrow 0$, the prototypes are distributed as uniformly as possible over the traits of every feature, yielding the largest possible separation between the prototypes in cultural space (which only coincides with the maximum of 1 allowed by Eq. (2.5) if the number of traits q is larger or equal to the number of prototypes k for every feature). This is achieved by a formulation in terms of the set of integer partitions I_k^q describing the possible ways of distributing the k prototypes over the q traits of a certain feature. The α parameter actually controls the probability distribution over the set I_k^q , via the “compactness” of the integer partitions in this set. Sec. 2.A.2 precisely describes how these probabilities are assigned and how the set I_k^q is computationally generated in the first place, for any combination of k and q . Once the prototypes are chosen, everything else is conditional on them, for both models.

According to the **Prototype Generation (PG)** model, each cultural vector is a partial realization of one of the prototypes. Each of the N cultural vectors is generated by copying a random sequence of traits from one of the k prototypes, while generating the other traits in a uniformly random way – choosing the prototype is done randomly for every vector. Then, a subset of the F features of length $\text{round}(\beta \cdot F)$ is randomly and independently selected for each vector and the traits of these features are copied from the prototype to the vector. Here, “round” returns the integer that is closest to its argument, while $\beta \in [0, 1]$ is a third model parameter, in addition to k and α (which are already needed for the purpose of specifying the prototypes, in the manner described above). The β parameter specifies the fraction of traits that are directly copied from the prototype, thus controlling for the expected distance between a vector and its prototype. The traits for the remaining features are generated randomly and independently, according to uniform feature-level probability distributions – the explicit random generation mentioned above. Thus, β also controls for the amount of explicitly random generation of traits. The PG method effectively specifies that there are k “classes” of cultural vectors and those of a certain class are located at a certain, β -controlled average distance from the associated cultural prototype. This is similar to the “Prototype Evolution” method of Ref. [13], although there are small differences in how exactly the vectors are generated in the two cases. Moreover, the method of Ref. [13] did not allow for controlling the expected cultural distance between the prototypes.

According to the **Mixed Prototype Generation (MPG)** model, each cultural vector is a combination of all prototypes, although an unbalanced combination, meaning that the numbers of traits copied from the different prototypes are deliberately unequal. The extent of this discrepancy is explicitly controlled via the third model parameter, which, like for PG, is called β . Although the exact

definition and usage of the $\beta \in (0, 1)$ parameter is different in MPG than in PG, its role is quite similar. Specifically, also in the context of MPG, β (indirectly) controls for the fraction of traits copied from the dominating prototype to the vector: more traits are copied from the dominating prototype if the discrepancy between the prototypes is higher. In addition to traits copied from the prototypes, there are traits that are generated in an explicitly random way, but in a small number. For each generated vector, this number is by construction not higher than the number of traits copied from the lowest-contributing prototype. Consequently, if there are k prototypes, the number of traits generated via explicit randomness does not exceed $F/(k+1)$. Thus, $1/(k+1)$ is an upper bound for the fraction of explicit randomness in an entire set of cultural vectors generated with MPG. It is also important to note that, like for PG, this fraction is controlled by β and that the upper bound is reached when β is in the limit of minimal imbalance. The limited usage of explicitly random trait generation by MPG means that cultural vectors are more strongly constrained by the prototypes, compared to PG. Still, MPG allows for generating a large variety of possible cultural vectors, since the k prototypes can mix in many different ways.

The MPG model needs a procedure of specifying, for each generated vector, the k values of the numbers of traits that are to be copied from the k prototypes, along with the number associated to explicitly random generation. These $k+1$ positive, integer numbers should add up to F and have their discrepancy controlled by the β parameter. Moreover, there is no reason to believe that the sequence of numbers associated to one β value should be the same across all generated vectors, so randomness should be involved in choosing these numbers. Therefore, the model needs a probabilistic way of drawing $k+1$ random, positive integers $\{t_1(\beta), \dots, t_{k+1}(\beta)\}$ satisfying $\sum_{l=1}^{k+1} t_l(\beta) = F$, such that their expected discrepancy is controlled via a single parameter β . The procedure chosen for this purpose is described below.

This procedure heavily relies on isometrically mapping the discrete $\{0, 1, \dots, F\}$ set of integers to the $[0, 1]$ interval of the real axis. For each generated vector, the latter interval is split into $k+1$ parts, by performing “cuts” in k randomly chosen points. In this manner, a sequence of $k+1$ preliminary weights $\{W_1, \dots, W_{k+1}\}$, subject to $\sum_{l=1}^{k+1} W_l = 1$ is numerically obtained. These weights are obviously independent of β and have a fixed expected discrepancy. A β -dependent transformation (explained below) is applied on the preliminary weights $\{W_1, \dots, W_{k+1}\}$, thus providing a sequence of β -dependent weights $\{w_1(\beta), \dots, w_{k+1}(\beta)\}$ satisfying $\sum_{l=1}^{k+1} w_l(\beta) = 1$, with expected discrepancy controlled by β . Finally, the sequence of β -dependent weights is converted back to the desired sequence $\{t_1(\beta), \dots, t_{k+1}(\beta)\}$. This final operation is non-trivial, requiring a self-consistent, joint rounding procedure, which is generally difficult to choose, since one cannot generally ensure that $w_l = \text{round}(t_l/F)$, $\forall l$ – a non-trivial problem of weight discretization. Here, a simple, pragmatic choice is made: converting the lowest k weights to the closest, lower integer, while converting the highest weight to the

integer needed for satisfying the summation constraint – this ensures that the highest weight, which should correspond to the dominating prototype, is converted to the highest integer.

The only aspect of MPG remaining to be explained is how the β -dependent weights $\{w_1(\beta), \dots, w_{k+1}(\beta)\}$ are obtained from the preliminary weights. This is done by raising the latter to a common power $p(\beta)$ and then normalizing:

$$w_l(\beta) = \frac{(W_l)^{p(\beta)}}{\sum_{l'=1}^{k+1} (W_{l'})^{p(\beta)}}, \quad (2.1)$$

where the common power $p(\beta) \in (0, +\infty)$ controls for the average discrepancy between these weights and maps to $\beta \in (0, 1)$ via:

$$p(\beta) = \tan\left(\beta \frac{\pi}{2}\right), \quad (2.2)$$

where the tangent is a convenient choice of a smooth, continuous function, with the appropriate domain and range. Thus, a value $\beta > 0.5$ implies a value $p > 1$ and a higher discrepancy of $\{w_1^p, \dots, w_{k+1}^p\}$ than that of $\{W_1, \dots, W_{k+1}\}$, while a value $\beta < 0.5$ implies a value $p < 1$ and a lower discrepancy of $\{w_1^p, \dots, w_{k+1}^p\}$ than that of $\{W_1, \dots, W_{k+1}\}$.

Before describing the fitting and the outcomes of the PG and MPG models, it is worth summarizing a few important aspects. Both models rely on the notion of cultural prototypes, which is currently formalized in a simplistic manner, which is only sensible for cultural spaces defined exclusively in terms of nominal features. The procedure for generating the prototypes is the same for both models and relies on two parameters, k and α , which specify, respectively, the number of prototypes and the expected distance between them. The differences between PG and MPG consist in how the cultural vectors are generated conditionally on the prototypes: for PG, every vector is in part a copy from one of the prototypes and in part explicitly random; for MPG, every vector is an imbalanced mixture of all prototypes and explicitly random to a much lower extent, which is how the “multiple-self” ingredient is implemented. Nonetheless, in both cases, there is a third model parameter, β , which governs, in different ways, the lengths of the randomly selected subsets of features whose traits that are copied from the prototypes. In both cases, β effectively controls for the expected distance between a vector and its (dominating) prototype, as well as for the fraction of explicit randomness.

2.3 Model fitting

Before applying the LTCD-STCB analysis on SCVs generated with either the PG or MPG models, it is useful to somehow constrain some of the free model parameters. This is done in terms of statistical quantities simpler than the LTCD

and the STCB measures, that can be evaluated on both empirical SCVs and on the model SCVs. On the empirical side, the quantities are averaged over several, empirical SCVs constructed by randomly selecting $N = 500$ cultural vectors from the 13000 available ones in Eurobarometer data set [34], while restricting to the nominal features – let “(EBM_n)” stand for the nominal part of the Eurobarometer data set. The empirical data is formatted according to the procedure explained in Ref. [14] (Chap. 1). On the model side, these quantities are averaged over many SCVs, of the same size N , that are realizable in the cultural space of (EBM_n), for the given combination of parameters – the prototypes are independently generated upon creating every model SCV.

The two simple quantities in terms of which the models are tuned to empirical data are the average and the standard deviation of the inter-vector distances in the SCV, which are here denoted by “AIVD” and “SIVD” respectively:

$$\text{AIVD} = \frac{2}{N(N-1)} \sum_{i < j} d_{ij}, \quad (2.3)$$

$$\text{SIVD} = \sqrt{\frac{2}{N(N-1)} \sum_{i < j} (d_{ij} - \text{AIVD})^2}, \quad (2.4)$$

where N is the number of cultural vectors and d_{ij} is the cultural distance, as defined and used in Refs. [14, 13, 12] (and Chap. 1). The notation $i < j$ denotes that the respective summation is carried out over all distinct pairs (i, j) . In the case of a fully-nominal cultural space, with which this study is dealing, d_{ij} reduces to the Hamming distance between the two sequences of symbols encoding cultural vectors i and j :

$$d_{ij} = 1 - \frac{1}{F} \sum_{l=1}^F \delta(x_i^l, x_j^l) = \frac{1}{F} \sum_{l=1}^F d_{ij}^l, \quad (2.5)$$

with, d_{ij} taking values within the $[0, 1]$ interval. Here, l iterates over the F nominal features, x_i^l, x_j^l are the traits of vectors i and j with respect to feature l and δ stands for the Kronecker-Delta function. The second equality shows that the cultural distance can be expressed as an average over feature-level contributions, which becomes useful below. Previous work has shown that an empirical SCV is characterized by a lower AIVD than its random counterpart and a higher SIVD than both its random and shuffled counterparts [12, 13]. The AIVD and SIVD quantities, which incorporate pairwise distance information, are conceptually different than what is often used in the context of cultural dynamics and of the Axelrod model, namely the size of the largest connected component, which can be regarded as an overall measure of similarity. Instead, the latter is somewhat similar to the STCB quantity explained and used in Sec. 2.4.

It is instructive to see that the expressions of AIVD and SIVD can be rewritten

in the following way:

$$\text{AIVD} = \frac{1}{F} \sum_{l=1}^F \frac{2}{N(N-1)} \sum_{i < j} d_{ij}^l, \quad (2.6)$$

$$\text{SIVD} = \sqrt{\frac{1}{F^2} \sum_{l, l'=1}^F \left[\frac{2}{N(N-1)} \sum_{i < j} d_{ij}^l d_{ij}^{l'} - \frac{4}{N^2(N-1)^2} \sum_{i < j} d_{ij}^l \sum_{i' < j'} d_{i'j'}^{l'} \right]}, \quad (2.7)$$

by using a feature-level cultural distance d_{ij}^l introduced via Eq. (2.5) – the transition from (2.4) to (2.7) was suggested by the SI of Ref [12].

Note that the AIVD can be understood as an average over feature-level AIVD contributions, which are represented by the expression within the l -summation of Eq. (2.6). It can be checked that the (nominal) feature-level AIVD contribution is a measure of how uniformly the N vectors are distributed over the possible traits of that feature. This is more obvious when expressing the expected value of the AIVD contribution in terms of probabilities associated to the traits, which is shown in Eq. (2.8) below. Thus, for an empirical SCV containing only nominal features, the AIVD is a measure of average uniformity of the empirical frequency distributions associated to the features. Consequently, the AIVD is also a measure of how subjective the questions/topics associated to the features are on average – when the frequencies of possible answers are more similar to each other, there is less justification to talk about a “better”, a “more correct” or a “more agreed upon” answer, so the question is inherently more subjective.

Also note that, in Eq. (2.7), the quantity inside the average over pairs of features (k, l) is the covariance between features k and l , defined in terms of the feature-level cultural distances. Given that this quantity is averaged over all possible pairs of features and that the square-root is a monotonous function, the SIVD encodes information about the pairwise correlations between features, although in a somewhat indirect way.

For both models, the choice made here is that of:

- tuning the α parameter in terms of the AIVD quantity (Eqs. (2.3), (2.6)), for any combination of values of the β and k parameters;
- tuning the β parameter in terms of the SIVD quantity (Eqs. (2.4), (2.7)), for any value of the k parameter, based on the previous fitting of α in terms of AIVD;
- simply repeating the tuning (and the LTCD-STCB analysis in Sec. 2.4) for several values of k .

This implies that, for every value of k , the tuning (or fitting) is done at two levels: the α -AIVD level and the β -SIVD level, the former being nested into the

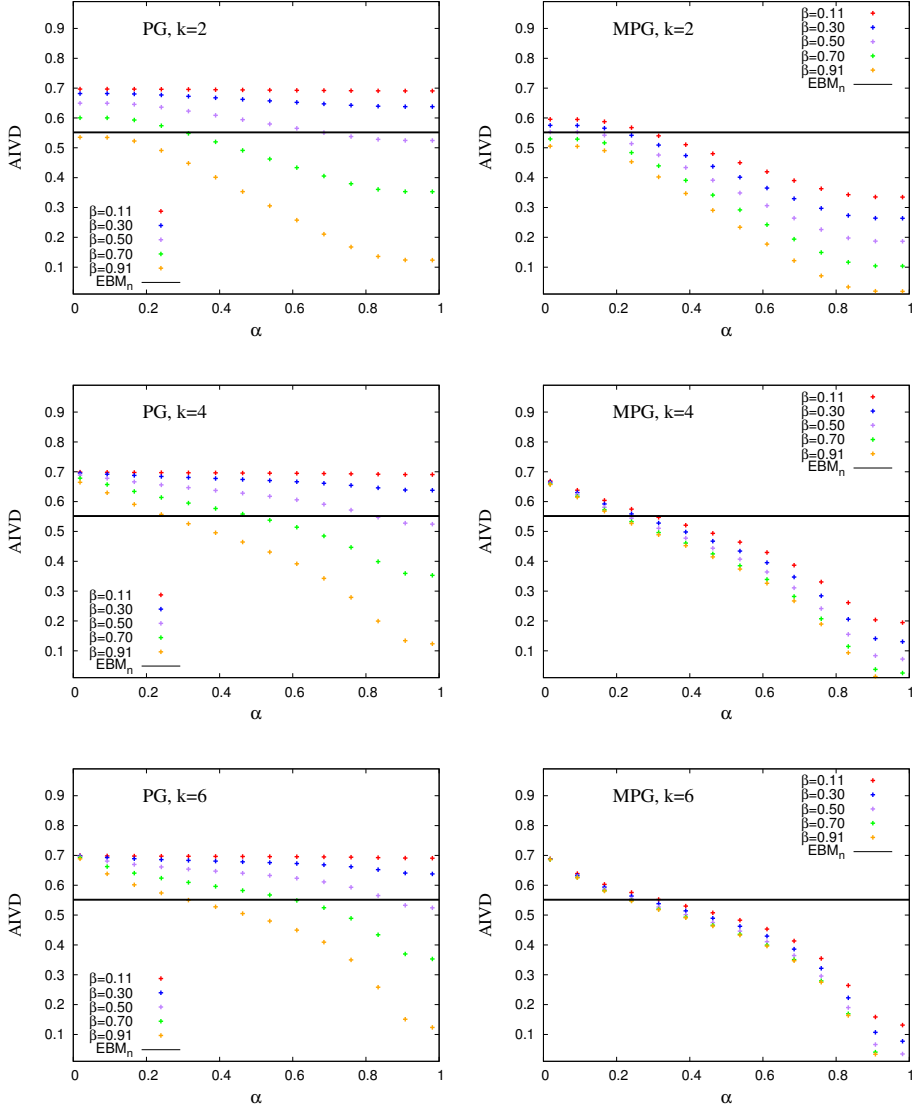


Figure 2.2: Dependence on model AIVD on the α parameter, for several values of the β parameter (legend), for $k = 2$ (top), $k = 4$ (center) and $k = 6$ (bottom) prototypes, for the PG (left) and MPG (right) models. The horizontal lines show the empirical AIVD uncertainty range (one standard error on each side of the mean).

latter. In practice, the fitting is carried out automatically, using a nested, 2-levels algorithm that relies on a modified bisection-type method for each level. The algorithm is precisely described in Sec. 2.C. In order to work, this approach relies on the assumption that there is one, unique solution for the fitting problem, for every value of k . This uniqueness is demonstrated via Figs. 2.2 and 2.3, which are also used for providing a general intuition of how the fitting works and of how the AIVD and SIVD quantities depend on α , β and k , for the two models.

Before entering this description, it is worth mentioning that the computer time for the fitting algorithm is greatly reduced by being able to evaluate the average (model) AIVD quantity analytically, in a manner that properly accounts for all SCVs that can be generated for any combination of k , α and β . While the calculation is described in detail in Sec. 2.B, a schematic understanding can already be provided here. The essential ingredient of the calculation is a simple, exact formula for the expected AIVD contribution of one feature of range q :

$$\langle \text{AIVD}(\{p_1, \dots, p_q\}) \rangle = 1 - \sum_{i=1}^q p_i^2, \quad (2.8)$$

which assumes that the probabilities of its traits $\{p_1, \dots, p_q\}$ are all known – see Sec. 2.B for the proof. For a discrete probability distribution, Eq. (2.8) is a measure of uniformity very similar to the Shannon entropy. Conditional on a specific choice of the prototypes, this set of probabilities (thus the feature-level probability distribution) is fully determined by the integer partition describing how the prototypes are distributed over the traits and by the fraction of traits that are randomly generated, the latter being controlled by β . In this context, Eq. (2.8) already assumes that an averaging is performed over SCVs generated from the same set of prototypes. One still needs to perform an average of this expression over integer partitions (Eq. (2.20) of Appendix Sec. 2.B), according to the probability distribution controlled by α (Eqs. (2.12) and (2.13) of Appendix Sec. 2.A.1), followed by another average over all features (Eq. (2.19) of Appendix Sec. 2.B), since different features will in general have different ranges q . At a superficial inspection, using a similar approach for analytically computing the SIVD quantity appears very complicated, if at all possible. Numerical calculations are instead employed for computing the (model) SIVD.

Fig. 2.2 deals with the first-level fitting. It shows the dependence of the analytically computed AIVD quantity (see above) on the α parameter, for several β values, for several k values and for both the PG and MPG models. Moreover, it shows the empirical AIVD uncertainty range² via the horizontal bands in the six panels. Thus, a solution of the first-level fitting is indicated by an intersection between a model curve of a given combination of k and β and the horizontal band. Note that, for either of the two models and for any combination of k and β , if a solution exists, this solution is actually unique. In order to understand

²An uncertainty range, as defined in Sec. 2.C, is the interval spanned by one standard mean error on each side of the mean.

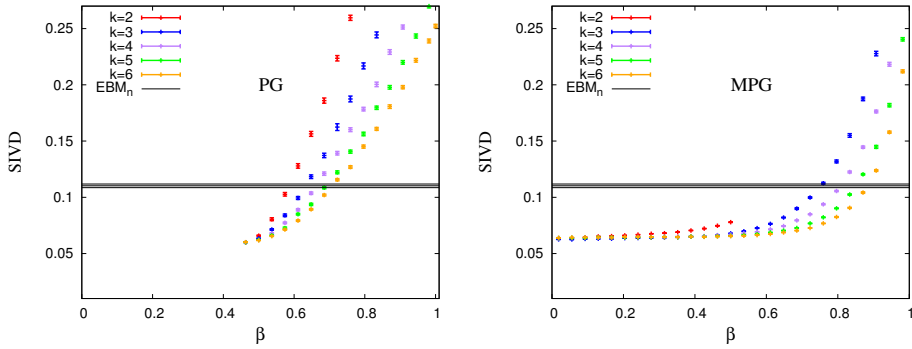


Figure 2.3: Dependence of model SIVD on the β parameter, for several values of the number of prototypes k (legend), for the PG (left) and MPG (right) models, where the α parameter is tuned such that the empirical AIVD is reproduced. The error bars of the points show the numerical uncertainty ranges, while the horizontal lines show the empirical SIVD uncertainty range (one standard error on each side of the mean).

the behavior implicit in Fig. 2.2, which is explained below, one should keep in mind that AIVD measures the average uniformity of the feature-level probability distributions.

First, it is worth focusing on the AIVD dependence on the α and β parameters. Note, on one hand, that for a given combination of k and β , the AIVD generally decreases with α , or at least remains constant. This is due to the fact that the AIVD decreases with decreasing distance between prototypes, thus with increasing α . For PG, this decrease is stronger for higher β values, since for low β value the uniformity is anyway high, because of the large fraction of randomly generated traits. For MPG, this β -dependence of the decrease is not that strong, since the fraction of randomly generated traits cannot exceed $1/(k+1)$. On the other hand, for a given combination of k and α , the AIVD generally decreases with increasing β . This is due to the fact that the AIVD decreases with decreasing fraction of randomly generated traits, thus with increasing β .

Second, it is worth focusing on the AIVD dependence on the number of prototypes k . For PG, for a given α , a larger number of prototypes k implies a higher AIVD, since traits copied from prototypes are more uniformly distributed, but this has a significant effect only for large β values, again due to the uniformity being anyway in place for small β values. For MPG, the corresponding behavior is more subtle. While for large, $\beta \rightarrow 1$ values, the AIVD still increases with increasing k at a given α (for the same reason as for PG), the AIVD(α) curves corresponding to small β approach the AIVD(α) curve corresponding to large $\beta \rightarrow 1$ with increasing k , rather than remaining in place (which is the case for

PG). This is related to the fact that the upper bound on the fraction of randomly generated traits $1/(k+1)$ decreases with increasing k , thus decreasing the role of β in controlling the AIVD via the uniform component of the feature-level probability distributions.

Fig. 2.3 deals with the second-level fitting. Everything shown in this figure relies on α already being tuned (at the first level) such that the empirical AIVD is matched – as apparent from Fig. 2.2, the tuned α value depends on β and on k . Fig. 2.3 shows the dependence of the numerically computed SIVD quantity (with uncertainty ranges) on the β parameter, for several k values and for both the PG and MPG models. Moreover, it shows the empirical SIVD uncertainty range via the horizontal bands in the two panels. Thus, a solution of the second-level fitting is indicated by an intersection between a model curve of a given k and the horizontal band. Note, again, that for either of the models and either of the k values, if a solution exists, this solution is actually unique. The exact technical procedure employed for producing any of the model points in Fig. 2.3 is described at the end of Sec. 2.C, followed by the explanation of the final choice of values for the α and β parameters, for use in the analysis of Sec. 2.4.

Note that the SIVD increases with β for both models and for all k values, suggesting that the extent of feature-feature correlation increases with decreasing distance between vectors dominated by the same prototype. For PG, all $\text{SIVD}(\beta)$ curves meet for some $\beta \approx 0.45$, at which point they also end. No points are plotted for lower β because α cannot be tuned in terms of AIVD, which can be understood from Fig. 2.2 when noticing the $\text{AIVD}(\alpha)$ curves of low β that do not cross the empirical line. For MPG, the $\text{SIVD}(\beta)$ curve of $k = 2$ ends at a value of $\beta \approx 0.5$, before crossing the empirical line, meaning that the MPG model cannot be entirely tuned when only 2 prototypes are used. No points are plotted for higher β because α cannot be tuned in terms of AIVD, which can be understood from Fig. 2.2, by noticing the $\text{AIVD}(\alpha)$ curves of $k = 2$ and high β that do not cross the empirical line. This is due to certain limitations of the current modeling paradigm, which are further discussed in Sec. 2.5.

2.4 Model Outcomes

Here, the most important results of this work are presented. The focus is on the LTCD-STCB analysis, applied to sets of cultural vectors generated with the PG and MPG models. The aim is to assess how well the two models reproduce the universal empirical patterns described in Ref. [14] (Chap. 1). Fig. 2.4 illustrates the results obtained with the two models, whereas Fig. 2.5 summarizes, for comparison purposes, the empirical results, focusing on the nominal part of the Eurobarometer dataset (EBM_n) – formatted according to the procedure explained in Ref. [14] (Chap. 1).

Before describing the results, it is worth recalling the main ingredients of the LTCD-STCB analysis. This is essentially a two-dimensional plot showing

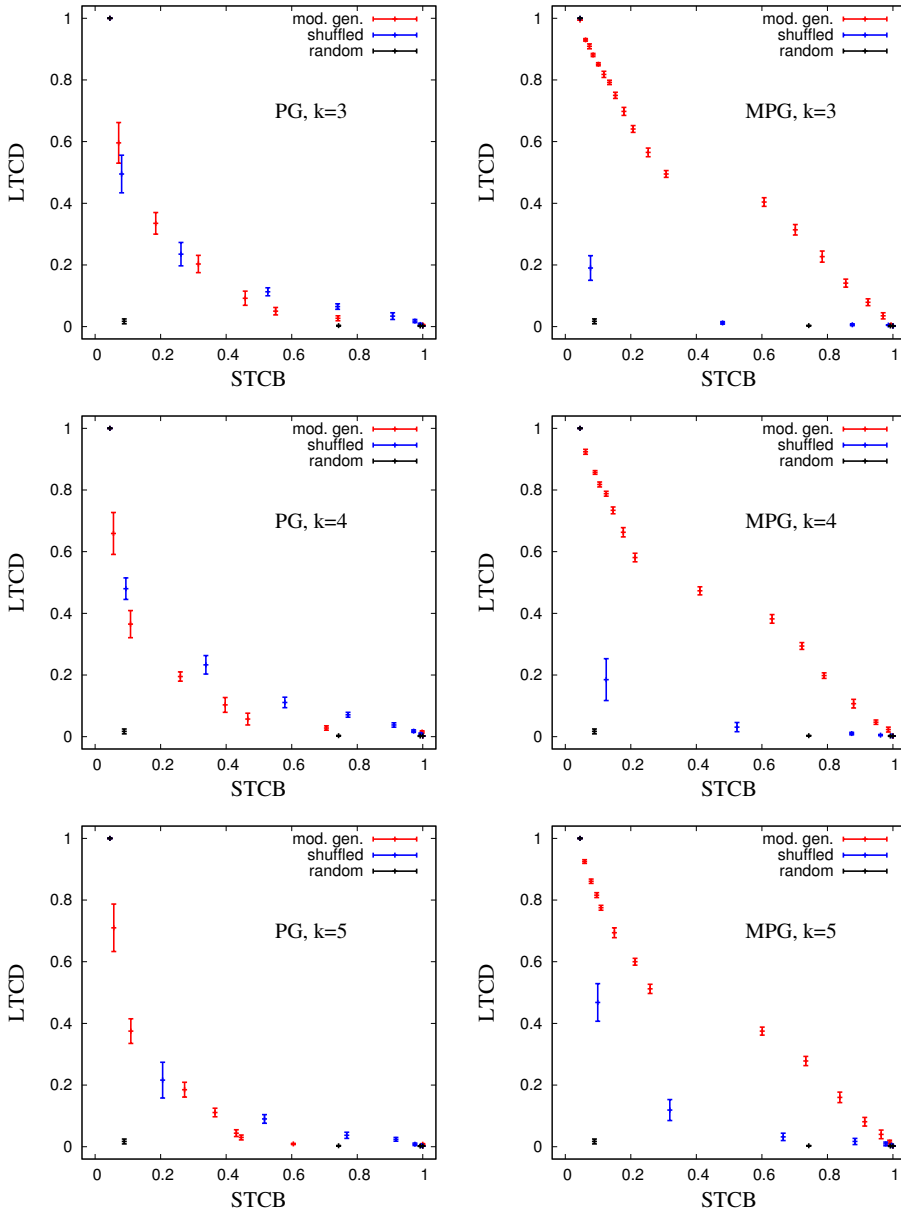


Figure 2.4: The correspondence between Long-term cultural diversity (LTCD) and short-term collective behavior (STCB) for a model-generated (red), a shuffled (blue) and a random (black) SCV obtained via the PG model (left) and MPG model (right), for $k = 3$ (top), $k = 4$ (centre) and $k = 5$ (bottom) prototypes. Error bars denote standard deviations over multiple trait dynamics runs. There are $N = 500$ elements in each set of cultural vectors.

the correspondence between the LTCD quantity vs the STCB quantity, both of them being evaluated on empirical, on shuffled and on random SCVs. Drawing the LTCD-STCB correspondence is made possible by the fact that, for each of the three scenarios, both quantities depend on the bounded-confidence threshold ω , which controls the maximal cultural distance over which social influence can act. On one hand, the LTCD quantity is a measure of cultural diversity after a long-term process of cultural dynamics driven by ω -bounded social influence, starting from an initial cultural state specified by the respective SCV. Essentially, it counts the number of distinct points in cultural space (commonly referred to as “cultural domains”) towards which the agents converge in the final state of a minimalistic, bounded-confidence Axelrod model. The STCB quantity is a measure of collective behavior (or social coordination) after a short-term process of opinion dynamics driven by ω -bounded social influence. Essentially, it is the standard deviation of the aggregate opinion distribution of the agent population, resulting from a minimalistic Cont-Bouchaud-type model applied to the (cultural) graph obtained by drawing a link for each pair of agents separated by a cultural distance smaller than ω . Mathematically, the two quantities, as functions of the bounded-confidence threshold ω , are captured by the following two expressions:

$$\text{LTCD}(\omega) = \frac{\langle N_D \rangle_\omega}{N}, \quad \text{STCB}(\omega) = \sqrt{\sum_A \left(\frac{S_A}{N} \right)_\omega^2}, \quad (2.9)$$

where N_D is the number of cultural domains in the final state of the Axelrod-type model, N is the number of agents (and cultural vectors) and S_A is the size of the A 'th of connected components in the ω -determined cultural graph. The average in the LTCD formula is taken over multiple simulations of the Axelrod-type model. The STCB quantity is calculated analytically, once the cultural connected components are found, based on the assumption of independent opinion-agreement within each connected component. An essential difference between the two quantities, reflected in the long-term/short-term distinction, consists of an idealized separation between two time-scales, in terms of the role that the SCV specified as input plays: cultural vectors, together with the distances between them, are assumed to be dynamical by the LTCD definition and static by the STCB definition, such that one deals with dynamics of vectors and with dynamics on vectors in the two cases respectively. The interested reader is referred to Refs. [14, 12] (and Chap. 1) for more details and remarks about the LTCD-STCB analysis.

For both the PG and the MPG models, the α and β parameters are tuned in the manner described in Sec. 2.3 for every value of the number of prototypes k , while the latter is simply iterated over. In Fig. 2.4, the LTCD-STCB plot is shown for the values $k = 3$, $k = 4$ and $k = 5$, for the PG (left) and the MPG (right) models. The value $k = 2$ is omitted since the α and β parameters could not be both tuned for MPG with two prototypes. All SCVs are generated using the cultural space of EBM_n , whose empirical SCVs also served for providing the AIVD and SIVD values in terms of which the tuning was conducted (Sec. 2.3).

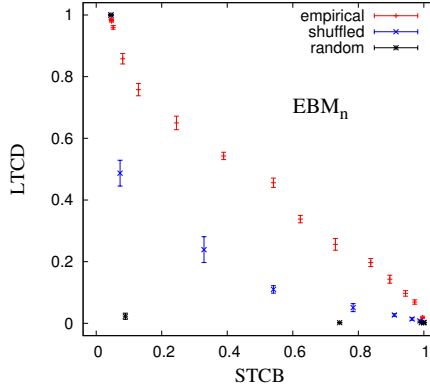


Figure 2.5: The correspondence between long-term cultural diversity (LTCD) and short-term collective behavior (STCB) for the empirical (red), shuffled (blue) and random (black) sets of cultural vectors, for the nominal part of the Eurobarometer data set (EBM_n). Error bars denote standard deviations over multiple cultural dynamics runs. There are $N = 500$ elements in each set of cultural vectors

When looking at Fig. 2.4, one should ask whether the universal, empirical patterns are reproduced by any of the six illustrated model scenarios. Qualitatively, the patterns are defined first in terms of a higher compatibility between LTCD and STCB in the model-generated SCV than in the shuffled SCV and a higher compatibility in the shuffled SCV than in the random one, second in terms of the model-generated LTCD-STCB curve being close to the second diagonal. These empirical features are visible in Fig. 2.5. It is clear that PG does not satisfy these criteria for any value of k . Indeed, the model-generated curve is far below the second diagonal for most of the relevant interval and often below the shuffled curve. MPG, however, appears to satisfy all these criteria for all k values, although for $k = 3$ it is not obvious that the shuffled curve is indeed above the random one, due to the lack of points in the lower-left corner. This has to do with the effective discreteness of the bounded-confidence threshold ω spectrum, due to the finite number of nominal features available – in other words, it is meaningless to split the ω axis into intervals that are smaller than the nearest-neighbor spacing of the cultural space lattice. For a direct comparison with analogous empirical curves, one should use Fig. 2.5, which shows the results of the LTCD-STCB analysis applied to EBM_n data. However, it is only meaningful to compare the qualitative nature of the empirical and the model curves, rather than the exact values, since, as discussed in Sec. 2.5, neither model has a maximum-likelihood nature, due to a certain simplicity in the way prototypes are formalized and chosen here. Still, MPG apparently does generate SCVs that are structurally similar to the empirical ones. Thus, the notion of cultural prototypes, even if implemented in a simplistic

way, can be used to reproduce the important, universal properties of empirical cultural states, as long as mixing of prototypes is in place.

2.5 Discussion

The purpose of this study was to develop a way of generating cultural states that reproduce the apparently universal properties of the empirical ones, namely those described by Ref. [14] (Chap. 1). This naturally calls for input from social science, in particular from social science theories that are intended to describe universal aspects of culture and society. There is an entire “class” of social science theories that appear relevant for this purpose, originating from either psychology or cultural anthropology [16, 17, 18, 19, 20], some of them being explicit attempts at unifying social science. All of them make use of cultural prototypes, although in somewhat different ways, under different names and numbers. Moreover, they had all been overlooked by previous studies of cultural dynamics, on which Ref. [14] (Chap. 1) largely builds: Ref. [13] was the first study that connected quantitative studies of cultural dynamics with these theories, via the generic, formal notion of cultural prototypes. For creating an instructive and compact context, this work focused on one of these theories, namely on Plural Rationality Theory (PRT).

There are several aspects justifying the focus on Plural Rationality Theory. First, its informal notion of cultural bias matches very well the more formal notion of cultural prototype, in the manner used in Ref. [13] and here. Second, it is more appealing from a natural science perspective than the others, in particular from a physics and complex systems perspective. This is largely due to various concepts that are qualitatively (and sometimes just implicitly) invoked by PRT, such as: energy landscapes, symmetry breaking, graph/network theory, dynamical systems, crossovers (possibly phase transitions), self-organization and fractals. Third, it explicitly claims to provide some insight into how preferences form: preferences are formed in the process of building social relations, while different patterns of relations (and types of institutional settings) go along with different conglomerates of preferences (the cultural biases). Finally, this dualism between patterns of relations on one hand and cultural biases on the other hand comes along with distinguishing between a “social plane” and a “cultural plane” of interacting human systems, while acknowledging the dynamical nature of both, as well as the strong coupling and interdependency between the two. Thus, PRT seems to resonate well, on one hand to research on social network structure and dynamics, on the other hand to research on cultural structure and dynamics.

Up to now, little work has been done to explore either of these two connections. While Ref. [13] and the present work are the first steps in exploring the latter connection, some steps have also been taken in exploring the former connection [35, 36]. Note, however, that Ref. [13] refers to several theories similar to PRT, without explicitly mentioning PRT, that Ref. [36] focuses on a social theory similar to PRT, while still discussing a connection with PRT and that Ref. [35]

works with an earlier, more rudimentary version of PRT, which gave less importance to the notions of “way of life”, “rationality” and “cultural bias”. Although the coupling between social dynamics and cultural dynamics is recognized and studied by quantitative complex systems research (for instance Refs. [9, 37]), this has been carried out in isolation from PRT.

In loose terms, each rationality of PRT has, as a “projection” on the cultural plane, one distinct cultural bias. These cultural biases correspond to the cultural prototypes used in this study. In agreement with Ref. [13], a cultural prototype is a combination of cultural traits, thus one point in cultural space – the limitations of this assumptions are extensively discussed below. Relying on these notions, two stochastic, structural models of culture are developed and studied here: Prototype Generation (PG) and Mixed Prototype Generation (MPG). It is important that, regardless of which model is used, once the prototypes and the remaining free parameters (parameter β , for either PG or MPG) are specified, one implicitly defines a cultural space distribution (CSD): a probability mass function taking the cultural space as a support, as defined in Ref. [14] (Chap. 1). Generating a set of N cultural vectors is then equivalent to selecting N points at random according to this distribution. Thus, the resulting cultural states are generated in a non-uniformly random way, with non-uniformities depending on the prototypes and on other model specifications.

For this study, the usage of both stochastic models is restricted to cultural spaces constructed only from sets of nominal features. This is due to the assumption that every prototype picks one and only one trait in any feature, which from a PRT perspective means that, upon answering a question under the influence of one cultural bias, a respondent can only provide one specific answer. In reality, even a specific cultural bias would generally point towards several answers, although with different probabilities, so it would be more realistic to say that every prototype corresponds to one probability distribution defined over that feature. Not allowing for this freedom makes this modeling paradigm incompatible to ordinal features, whose associated traits are by construction sorted along an axis, in which case it is not reasonable to assume that a prototype points to one trait of a feature with full probability and to its nearest-neighbors with zero probability. Nonetheless, the paradigm is reasonably compatible with nominal features, in which case the distance between any two traits of one feature is anyway assumed to be the same.

The current study belongs to a preliminary, simplistic paradigm which makes use of what one may call “sharp prototypes”. A more realistic paradigm, which would account for the probabilistic nature of the cultural biases, would make use of what one may call “diffuse prototypes”. Using sharp prototypes comes at the cost of not having enough flexibility to reproduce the empirical, feature-level frequency distributions, with either of the two models, since every prototype corresponds to a probability distribution entirely peaked on one trait. Instead, using diffuse prototypes would allow this by enforcing, for every feature, that the empirical distribution is a linear combination of the prototype distributions. Nonetheless,

as shown in Sec. 2.3, both models are still able to reproduce the empirical average uniformity of the feature-level frequency distributions, namely the AIVD quantity. This is partly due to both models making some use of uniformly-random trait generation, independently of the prototypes. This translates to a flat noise component in the probability distribution of every feature, which in a sense compensates for the rigid peaks of the sharp prototypes. When also considering the results of Sec. 2.4, the usage of sharp prototypes restricted to nominal variables appears to be enough as a proof of concept. This justifies further research towards the more sophisticated paradigm relying on diffuse prototypes. Although this is left for future studies, it is worth contemplating upon, in order to better understand the purpose, greater context and limitations of the current paradigm.

Working with diffuse prototypes should go hand in hand with a method of inferring them from data. One can imagine doing this by applying a sensible clustering method on the empirical set of cultural vectors, followed by a sensible method of constructing one diffuse cultural prototype from every cluster, as a probabilistic entity that is representative of that cluster. The main advantage of this approach is that once the prototypes are constructed and provided as input to a sensible stochastic model, the artificial SCVs generated with this model would be close-to-representative of the same distribution in cultural space as the empirical SCV on which the method is applied in the first place. This means that the model would have a maximum-likelihood flavor, and could be used for generating synthetic data, which would also reproduce the feature-level frequency distributions.

By contrast, the approximation of sharp prototypes used here is too strong to be employed together with a method of inferring them from data. Instead, sharp prototypes are being assigned to randomly chosen positions in the given cultural space. On one hand, the fact that the prototypes are randomly chosen makes any model symmetric up to any permutation of the traits of any feature, as long as all features are nominal, which is the case here, a symmetry which is broken by an empirical SCV and also by an artificial SCV generated from a specific choice of the prototypes. On the other hand, the fact the prototypes are sharp does not allow for the exact frequency distribution of a specific feature to be reproduced, not even up to a permutation of the traits. Still, after parameter tuning, one should expect from a good model to provide a cultural space distribution whose rough “shape” is compatible with the empirical data, though the “orientation” and the structural details implied, for instance, by the feature-level distributions would not be compatible. This should reflect in roughly reproducing the universal LTCB-STCB patterns emphasized in Ref. [14] (Chap. 1): on one hand, the formulation of the LTCB and STCB observables is also symmetric up to permuting the traits of any feature, and thus independent of the “orientation”; on the other hand, the empirical, feature-level frequency distributions should heavily depend on the specific data set, thus being of little relevance for the universal patterns.

There are various aspects that make the random generation of prototypes sensible for the purpose of the present work. First, results are evaluated for

various values of the number of prototypes k , which is considered a free parameter for both the PG and MPG model. Second, the expected prototype-prototype distance is controlled for via parameter α . Third, for every choice of parameters, the prototypes are independently drawn for each realized cultural state in the set used for computing the model AIVD and SIVD quantities for fitting purposes. These compensate somewhat for not inferring the prototypes from empirical data.

In order to give an example of how the sharp prototypes approximation can be pushed beyond its limits, it is worth recalling that fitting the MPG model is not possible for $k = 2$ prototypes, as pointed out at the end of Sec. 2.3: the α parameter can be successfully tuned in terms of the AIVD only for small β values, which do not allow for the subsequent fitting of the β parameter in terms of the SIVD. This is related to there being at least $q = 3$ traits associated to every nominal feature selected from the Eurobarometer data set, while there are only two, prototype-induced peaks in the model probability distribution of every feature, on top of the uniform component. Since the integrated probability of the uniform component cannot exceed $1/k$ by construction, all the distributions are bound to be relatively non-uniform, such that the empirical average uniformity is only attained for small- α (few coincidences between the prototype-induced peaks) and small- β (large uniform component) combinations. This does not hold for the PG model, as in this case the integrated probability of the uniform component can attain any value between 0 and 1. Nonetheless, if $k > 2$, the fitting of the MPG leads to generated cultural states that reproduce much better the universal empirical patterns than PG. This justifies considering MPG the successful model, while emphasizing the importance of the mixing ingredient, which validates the multiple self assumption.

When thinking in terms of the feature-level probability distributions, it might seem that the MPG and PG models are not that different from each other. As mentioned above, for both models, if there are k prototypes, the probability distribution of a certain feature would consist of k peaks of equal probability contents and of a uniform component associated to the explicitly random trait generation. Although the probability content of the uniform component of MPG is bounded from above, that of PG is not bounded in any way, so one might think that MPG is just a particular realization of PG. However, this reasoning is misleading, as it focuses on partial information encoded in the feature-level probability distributions, disregarding the rest of the information encoded in the complete cultural space distribution. With PG, a cultural vector whose trait, with respect to a certain feature, is generated under the probability peak of a certain prototype will have its trait generated, with respect to another feature, under the well-determined probability peak of the same prototype or under the uniform component. By contrast, with MPG, a cultural vector whose trait, with respect to a certain feature, is generated under the probability peak of a certain prototype, will have its trait generated, with respect to another feature, under the probability peak of any prototype – though with a higher likelihood under the peak of the dominating prototype – or under the uniform component. Thus, for the same choice

of the prototypes and the same extent of explicitly random generation of traits (and consequently the same AIVD), PG implies a different level of cross-feature correlation and a different shape of the cultural space distribution than MPG. This conceptually explains the impact of the mixing ingredient.

Although this study does not attempt at providing a complete mathematical theory of trait dynamics and formation, one can argue that the MPG model qualifies as a good effective,³ static description of (generic snapshots of) trait dynamics. This static description is inspired by Plural Rationality Theory which, although originating in cultural anthropology, does seem to integrate notions of both psychology and of a (complex) systems based understanding of society. Although it is formulated in an a qualitative, informal way, Plural Rationality Theory and related research should be of use for developing a complete formal theory of trait dynamics, at least as a source of guidance and inspiration.

2.6 Summary and conclusions

This study was dedicated to developing and testing a stochastic model for generating cultural states that would be structurally similar to the empirical ones. The aim was to reproduce the universal, empirical properties pointed out in Ref. [14] (Chap. 1), while relying on some social science hypothesis. Following up on previous work, the idea of cultural prototypes was used for this purpose. The study first tested the hypothesis that each cultural vector is a partial realization of one prototype and random for the rest, which is what was previously assumed. This turned out to be insufficient for reproducing the empirical patterns. Instead, one has to assume that each cultural vector is a combination, or mixture of all prototypes, although still dominated by either of them, which is what the MPG model encodes. This additional, mixing ingredient is actually suggested by the same social science theories that inspired the prototypes idea in the first place. In this specific, social science context, this aspect is often referred to as “the multiple-self”. These results provide indirect evidence for social science theories like PRT, that postulate, in one way or another, some notion of cultural prototypes, along with some associated notion of mixing.

Still, there is a certain rigidity in the way prototypes are currently formalized (Sec. 2.5), related to the assumption that every prototype corresponds to one and only one value of every cultural variable, instead of corresponding to a probability distribution over the variable. This makes the cultural space distribution induced by the successful, MPG model generally incompatible with the cultural space frequency distribution with respect to which it is fitted. As it stands, MPG is far from being a maximum-likelihood type of model and thus cannot be used to generate synthetic data. Nonetheless, this is arguably achievable once diffuse

³Note that “effective description of” stands for “description of the effects of”, for “approximate description” or for “phenomenological description”, as used in the physics literature, rather than for “successful or “efficacious”.

prototypes are used instead of sharp ones, while being inferred from the data rather than randomly chosen. In this sense, this work can be seen as an important step towards a realistic, maximum-likelihood model of empirical cultural states, and towards generating synthetic sets of cultural vectors. Moreover, MPG can be considered an effective description of the outcome of trait dynamics, since the generated cultural states seem to reproduce the generic structure of the empirical ones. The LTCD-STCB analysis, used for validating this effective theory, could also be used for validating a more fundamental, dynamical theory of culture. It appears likely that Plural Rationality Theory has more to say for aiding the development of such a theory.

Appendices

2.A Controlling the generation of prototypes

This section describes the calculation of probabilities attached to sets of cultural prototypes employed by the PG and MPG models defined in Sec. 2.2. These probabilities are collectively controlled via a parameter (α), which effectively dictates the expectation value of the average prototype-prototype cultural distance for one set of prototypes. The assignment of traits to prototypes is conducted independently for every feature, so the discussion is reduced to assigning probabilities to prototype-to-trait mappings at the level of a single feature. Furthermore, since prototype generation neglects empirical occurrence frequencies of specific traits, the problem is symmetric with respect to permutations of the traits, so the discussion is further reduced to assigning probabilities to “topologies” of prototype-to-trait mappings at the level of a single feature. Mathematically, such a topology is an “integer partition”. Integer partitions turn out to be the mathematical objects to which elementary probabilities are to be assigned. Sec. 2.A.1 explains the procedure for assigning the probabilities to integer partitions, while Sec. 2.A.2 explains the procedure for generating the integer partitions.

2.A.1 Integer partition probabilities

Let I_k be the set of all integer partitions of k elements, where an integer partition of k elements is an ordered sequence of integers that add up to k , also called “parts”. Let the ordered sequence $(k_1, \dots, k_s) \in I_k$ be one generic element of this set, where s counts the number of non-zero parts. This notation implies that the parts are sorted for descending values $k_i \geq k_{i+1} \forall i \in \{1, \dots, s-1\}$ and that they add up to $k = \sum_{i=1}^s k_i$. For instance, $(3, 2, 2, 1)$ is an integer partition of 8 elements with 4 parts. For the purpose of this work, an element of the integer partition corresponds to one prototype. For a specific choice of the prototypes

and a specific feature, an integer partition is a representation of how the prototypes are distributed over the traits of this feature, up to a permutation of these traits. Thus, when the fraction of traits that are randomly generated vanishes, the probabilities of the traits are just the normalized part sizes – in the example above, the ordered sequence of probabilities associated to the traits would be $(\frac{3}{8}, \frac{2}{8}, \frac{2}{8}, \frac{1}{8})$. Random trait generation then simply introduces a uniform, noise component to the feature probability distribution, whose contribution increases with the fraction of traits that are randomly generated. Thus, the integer partition is in any case a proxy for the feature probability distribution, regardless of which stochastic model is used.

Let $c(k_1, \dots, k_s)$ be the “compactness” of integer partition $(k_1, \dots, k_s)$, defined by:

$$c(k_1, \dots, k_s) = \sum_{i=1}^s \frac{k_i(k_i - 1)}{2}, \quad (2.10)$$

which counts the number of pairs of elements belonging to the same part. For instance, the compactness of integer partition $(3, 2, 2, 1)$ is $c(3, 2, 2, 1) = 3^2 + 1^2 + 1^2 + 0^2 = 11$. The compactness thus counts the prototype-prototype coincidences for one feature. In light of the above paragraph, a small compactness implies a high uniformity for the feature probability distribution and thus a high value of the associated (feature-level) AIVD contribution.

Let I_k^q be the set of integer partitions of k elements of at most q parts (which implies that $I_k^q \subseteq I_k$). This definition is needed for working with features with range $q < k$. Furthermore, **let** $c_{k,q}^{\min}$ and $c_{k,q}^{\max}$ be the minimal and maximal compactness values attainable by the elements of I_k^q . These notions are needed for normalizing generic compactness values. They formally read:

$$\begin{aligned} c_{k,q}^{\min} &= c(\lambda') \cdot (\lambda' \in I_k^q \wedge \nexists \lambda \in I_k^q \cdot (c(\lambda) < c(\lambda'))), \\ c_{k,q}^{\max} &= c(\lambda') \cdot (\lambda' \in I_k^q \wedge \nexists \lambda \in I_k^q \cdot (c(\lambda) > c(\lambda'))), \end{aligned} \quad (2.11)$$

where the “.” (dot) notation stands for “with the property that”.

At this point, it is possible to define an non-normalized probability mass function parametrized by α over the discrete set of integer partitions I_k^q , function whose shape would depend on α . High α values correspond to integer partitions of high compactness values being favored over those of low compactness values, while low α values correspond to integer partitions of low compactness values being favored over those of high compactness values. For simplicity, the function is chosen to be monotonous when re-expressed in terms of compactness. A simple choice for such a function, denoted here by $\rho_{k,q}^\alpha$, is given by:

$$\rho_{k,q}^\alpha(\lambda) = \exp \left\{ \tan \left[(2\alpha - 1) \frac{\pi}{2} \right] \frac{2c(\lambda) - c_{k,q}^{\max} - c_{k,q}^{\min}}{c_{k,q}^{\max} - c_{k,q}^{\min}} \right\}, \quad (2.12)$$

where the inner fraction linearly maps the compactness $c(\lambda)$ from interval $[c_{k,q}^{\min}, c_{k,q}^{\max}]$ to interval $[-1, 1]$, while the argument of the tan function linearly maps α from interval $(0, 1)$ to interval $(-1, 1)$, from where it is further mapped to $(-\infty, \infty)$ by the tan function. In this manner, the function is increasing with $c(\lambda)$ for $\alpha > 0.5$ (implying a relatively low expectation value of average prototype-prototype separation), the function is decreasing with $c(\lambda)$ for $\alpha < 0.5$ (implying a relatively high expectation value of average prototype-prototype separation) and the function is a constant of $c(\lambda)$ for $\alpha = 0.5$. The actual probability $P_{k,q}^\alpha(\lambda)$ associated to integer partition λ can then be obtained via the normalization:

$$P_{k,q}^\alpha(\lambda) = \frac{\rho_{k,q}^\alpha(\lambda)}{\sum_{\lambda \in I_k^q} \rho_{k,q}^\alpha(\lambda)}, \quad (2.13)$$

with the sum in the denominator being taken over all integer partitions in I_k^q .

2.A.2 Integer partition generation

Let $I \stackrel{\text{d}}{=} \{0^I, 1^I\} \cup I_1 \cup I_2 \cup \dots$ be the set of all integer partitions of any size, together with a “null” element 0^I and a “unity” element 1^I , which are meaningful in relation to the $\oplus$ operation defined below and are needed for keeping some of the following definitions compact and self-consistent.

Let the integer partition “merging” $\oplus : I \times I \rightarrow I$, acting on two integer partitions of k_a and k_b elements, with s_a and s_b parts respectively, be defined in the following way:

$$(k_1^a, \dots, k_{s_a}^a) \oplus (k_1^b, \dots, k_{s_b}^b) = (k_1, \dots, k_s), \quad (2.14)$$

producing another integer partition of $k = k_a + k_b$ elements and $s = s_a + s_b$ parts, such that the sequence of parts in the resulting partition is a sorted merging of the two original sequences of parts. For instance: $(3, 2, 2, 1) \oplus (4, 2) = (4, 3, 2, 2, 2, 1)$. Moreover, any integer partition $\lambda \in I$ satisfies $\lambda \oplus 0^I = 0^I$ and $\lambda \oplus 1^I = \lambda$.

Let the integer partition “multi-merging” $\otimes : I \times \mathcal{P}(I) \rightarrow \mathcal{P}(I)$, where $\mathcal{P}(I)$ is the set of all subsets of I , be defined by:

$$\alpha \otimes \{\alpha_1, \dots, \alpha_\sigma\} = \{\alpha \oplus \alpha_1, \dots, \alpha \oplus \alpha_\sigma\}, \quad (2.15)$$

where $\alpha, \alpha_1, \dots, \alpha_\sigma \in I$ are all integer partitions. The $\otimes$ operation produces a set of integer partitions of σ elements from an initial set of integer partitions of the same size and another integer partition α , by merging α with each element α_i in the initial set via the $\oplus$ operation.

Relying on the notions above, the following recursive definition of function $\text{sip}(k, m_L, m_V) : \mathbb{N} \times \mathbb{N}^* \times \mathbb{N}^* \rightarrow \mathcal{P}(I)$ encodes the procedure for generating the set of integer partitions of k elements, of maximally m_L parts, with maximal part

value m_V :

$$\begin{aligned} \text{sip}(k, m_L, m_V) &= \\ &= \begin{cases} \{1^I\} & k = 0, \\ \{0^I\} & k > m_L \cdot m_V, \\ \{(m_V, \dots, m_V)_{m_L \text{ entries}}\} & k = m_L \cdot m_V, \\ \bigcup_{x \in \overline{1, \min(k, m_V)}} [(x) \otimes \text{sip}(k - x, m_L - 1, x)] & \text{else.} \end{cases} \end{aligned} \quad (2.16)$$

definition inspired by Ref. [38], where the order of the four cases matters, in the sense that one case is considered only if none of the conditions of the above cases is valid. The last line returns the set resulted from the reunion “ $\cup$ ” of all sets of integer partitions of type $(x) \otimes \text{sip}(k - x, m_L - 1, x)$, where x spans the indicated interval. This general formulation, which also takes the maximal part value m_V as argument, is required for a compact recursive definition. But of actual interest for this work is the set of integer partitions of k elements and maximal part value q , I_k^q , given by:

$$I_k^q = \text{sip}(k, q, k) - \{0^I, 1^I\}, \quad (2.17)$$

where the last part of the expression takes out the null and/or the unity element, which might be present in the set of integer partitions as leftovers from the computation. Here we explicitly show how the sip function works when calculating the set of integer partitions of 4 elements of maximally 3 parts, given by $I_4^3 = \text{sip}(4, 3, 4) - \{0^I, 1^I\}$, where:

$$\begin{aligned} \text{sip}(4, 3, 4) &= \bigcup_{x \in \overline{1, 4}} (x) \otimes \text{sip}(4 - x, 2, x) = \\ &= [(1) \otimes \text{sip}(3, 2, 1)] \cup [(2) \otimes \text{sip}(2, 2, 2)] \cup [(3) \otimes \text{sip}(1, 2, 3)] \cup [(4) \otimes \text{sip}(0, 2, 4)] = \\ &= [(1) \otimes \{0^I\}] \cup [(2) \otimes \bigcup_{x \in \overline{1, 2}} (x) \otimes \text{sip}(2 - x, 1, x)] \cup \\ &\cup [(3) \otimes (1) \otimes \text{sip}(0, 1, 1)] \cup [(4) \otimes \{1^I\}] = \\ &= \{0^I\} \cup [(2) \otimes [(1) \otimes \text{sip}(1, 1, 1)] \cup [(2) \otimes \text{sip}(0, 1, 2)]] \cup [(3) \otimes (1) \otimes \{1^I\}] \cup \{(4)\} = \\ &= \{0^I\} \cup [(2) \otimes [(1) \otimes \{(1)\}] \cup [(2) \otimes \{1^I\}]] \cup [(3) \otimes \{(1)\}] \cup \{(4)\} = \\ &= \{0^I\} \cup [(2) \otimes [\{(1, 1)\} \cup \{(2)\}]] \cup \{(3, 1)\} \cup \{(4)\} = \\ &= \{0^I\} \cup [(2) \otimes \{(1, 1), (2)\}] \cup \{(3, 1), (4)\} = \\ &= \{0^I\} \cup \{(2, 1, 1), (2, 2)\} \cup \{(3, 1), (4)\} = \\ &= \{0^I, (2, 1, 1), (2, 2), (3, 1), (4)\}, \end{aligned} \quad (2.18)$$

yealding $I_4^3 = \{(2, 1, 1), (2, 2), (3, 1), (4)\}$, which is the expected result.

2.B Analytic calculations of model average inter-vector distance

This section explains the analytic calculation for the expectation value of the average inter-vector distance (AIVD) for sets of cultural vectors generated using either the PG or MPG model. The first part of this section just gives the essential formulas – Eqs. (2.19) and (2.20)) are common for the two models; the difference between the model becomes apparent when comparing Eq. 2.21 with Eq. 2.22. The second part gives the proof Eq. (2.8), which is the basis for Eq. (2.20).

The expectation value of the AIVD, as a function of the three model parameters k, α, β is given by the average over the feature-level expectation values:

$$\langle \text{AIVD} \rangle_{\alpha, \beta}^k = \frac{1}{F} \sum_q n_q \langle \text{AIVD} \rangle_{\alpha, \beta}^{k, q}, \quad (2.19)$$

where the sum goes over all possible values ranges q and n_q is the number of features with range q , with $\sum_q n_q = F$ being implicitly satisfied, where F is the number of features. Note that the feature-level contribution also depends on q . In turn, this contribution is given by:

$$\begin{aligned} \langle \text{AIVD} \rangle_{\alpha, \beta}^{k, q} = 1 - \sum_{(k_1, \dots, k_s)}^{I_k^q} P_{k, q}^{\alpha}(k_1, \dots, k_s) \cdot \\ \cdot \left\{ \sum_{i=1}^s \left[\pi_{\beta, F}^k \frac{k_i}{k} + (1 - \pi_{\beta, F}^k) \frac{1}{q} \right]^2 + (q - s) \left(\frac{1 - \pi_{\beta, F}^k}{q} \right)^2 \right\}, \end{aligned} \quad (2.20)$$

which is essentially a weighted averaging of Eq. (2.8) over the set of integer partitions $(k_1, \dots, k_s) \in I_k^q$, where the weights are the integer partition probabilities $P_{k, q}^{\alpha}(k_1, \dots, k_s)$. These are calculated in the manner described in Sec. 2.A.1, while the integer partitions themselves are generated in the manner described in Sec. 2.A.2. The set of p_i 's of Eq. (2.8) depends on the integer partition in the manner illustrated between the braces of Eq. (2.20), where the first term accounts for the s traits that are covered by the (non-zero) elements of the integer partition, namely those under the peak(s) of one (or more) prototype and under the flat noise component, while the second term accounts for the remaining $q - s$ traits, namely those that are only under the flat noise component. The dependence on whether the PG or the MPG model is used is captured by $\pi_{\beta, F}$, which is the average fraction of traits directly copied from prototypes, given by:

$$\pi_{\beta, F}^k = \frac{\text{round}(\beta F)}{F}, \quad (2.21)$$

for PG, where the “round” function accounts for the fact that only integer numbers

of traits can be copied, and:

$$\pi_{\beta,F}^k = \frac{1}{|W_{k+1}^\beta|} \sum_w^{W_{k+1}^\beta} w, \quad (2.22)$$

for MPG, where w iterates over all values of W_{k+1}^β , which is a large sequence of lowest MPG discrete weights (see Sec. 2.2), which are numerically generated during a previous step, for each used combination of (k, β) values. $|W_{k+1}^\beta|$ is the number of elements in this sequence of discrete weights. For this study, $|W_{k+1}^\beta| = 10^5$ elements were generated for every (k, β) combination, which allows for a very precise numerical calculation of $\pi_{\beta,F}^k$ in the case of MPG.

The consistency between the analytical AIVD calculation explained above and the numerical calculation is illustrated here via Fig. 2.6. The expected AIVD value is shown as a function of the β parameter, for 5 values of the α parameter and 3 values of the k parameter, for both the PG and MPG models. The analytical values are shown by the lines, while the numerical ones are shown by the dots, which have small, almost indiscernible error bars attached. For the numerical case, 50 sets of $N = 500$ cultural vectors are generated for each combination of parameters. Note that the numerical profiles follow closely the analytical ones, with small deviations that are consistent with the expected fluctuations of the mean.

It is now worth presenting a proof of Eq. (2.8), on which Eq. (2.20) is based. Consider a feature with q traits and a set of a-priori probabilities $\{p_1, \dots, p_q\}$ attached to them. Then, the entry of each cultural vector generated with respect to this feature is an independent, random choice from the q traits, according to the probability mass function $(p_1, \dots, p_q)$. Thus, the expected AIVD contribution from N cultural vectors is given by:

$$\begin{aligned} \langle \text{AIVD}(\{p_1, \dots, p_q\}) \rangle &= 1 - \frac{2}{N(N-1)}. \\ \sum_{x_1, \dots, x_q}^{x_1 + \dots + x_q = N} \sum_{i=1}^q \frac{x_i(x_i - 1)}{2} f\left(N, \frac{x_1, \dots, x_q}{p_1, \dots, p_q}\right) &= 1 - \frac{2}{N(N-1)}. \\ \sum_{i=1}^q \sum_{x_i}^{x_i \leq N} \frac{x_i(x_i - 1)}{2} \sum_{x_1, \dots, x_{i-1}, x_{i+1}, \dots, x_q}^{x_1 + \dots + x_{i-1} + x_{i+1} + \dots + x_q = N - x_i} f\left(N, \frac{x_1, \dots, x_q}{p_1, \dots, p_q}\right) &= \\ &= 1 - \frac{2}{N(N-1)} \sum_{i=1}^q \sum_{x_i}^{x_i \leq N} \frac{x_i(x_i - 1)}{2} S_i, \end{aligned} \quad (2.23)$$

where $f\left(N, \frac{x_1, \dots, x_q}{p_1, \dots, p_q}\right)$ denotes the probability that the N independent, random variables fill the q traits with the frequency distribution $(x_1, \dots, x_q)$, given the associated probability distribution $(p_1, \dots, p_q)$, where $\sum_{i=1}^q x_i = N$. This is conventionally called the multinomial distribution. In the above derivation, S_i stands

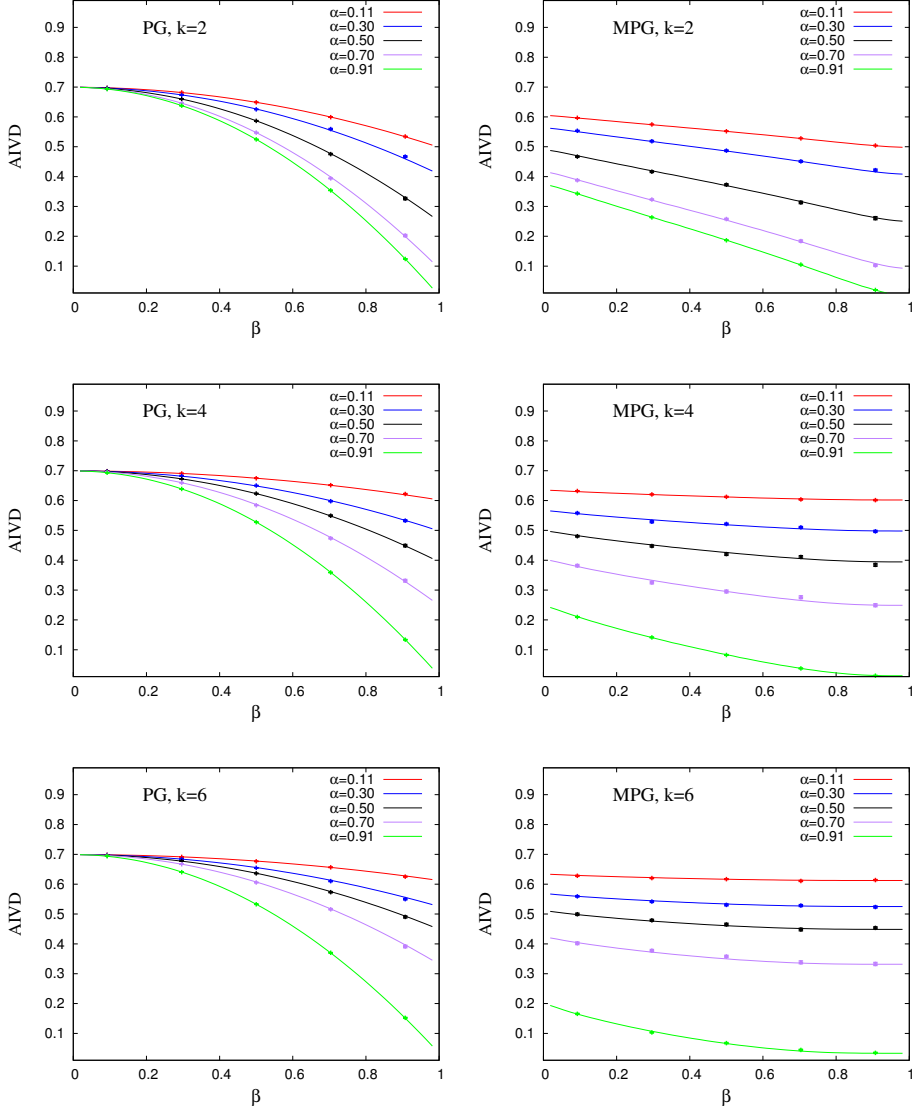


Figure 2.6: Comparison between numerical (dots) and analytical (line) expected AIVD as a function of β , for the PG (left) and MPG (right) models, with $k = 2$ (top), $k = 4$ (center) and $k = 6$ (bottom) prototypes, for several values of α (legend).

for the summation over all elements of the multinomial except that which has a certain, x_i number of entries for the i th trait, which can be further manipulated:

$$\begin{aligned}
 S_i &= \sum_{\substack{x_1, \dots, x_{i-1}, x_{i+1}, \dots, x_q \\ x_1 + \dots + x_{i-1} + x_{i+1} + \dots + x_q = N - x_i}} f\left(N, \begin{smallmatrix} x_1, \dots, x_q \\ p_1, \dots, p_q \end{smallmatrix}\right) = \sum_{\substack{x_1, \dots, x_{i-1}, x_{i+1}, \dots, x_q \\ x_1 + \dots + x_{i-1} + x_{i+1} + \dots + x_q = N - x_i}} \frac{N!}{x_1! \dots x_{i-1}! x_{i+1}! \dots x_q!} p_1^{x_1} \dots p_{i-1}^{x_{i-1}} p_i^{x_i} p_{i+1}^{x_{i+1}} \dots p_q^{x_q} \\
 &= p_i^{x_i} \frac{N!}{x_i! (N - x_i)!} \cdot \sum_{\substack{x_1, \dots, x_{i-1}, x_{i+1}, \dots, x_q \\ x_1 + \dots + x_{i-1} + x_{i+1} + \dots + x_q = N - x_i}} \frac{(N - x_i)!}{x_1! \dots x_{i-1}! x_{i+1}! \dots x_q!} p_1^{x_1} \dots p_{i-1}^{x_{i-1}} p_{i+1}^{x_{i+1}} \dots p_q^{x_q} = \\
 &= \binom{N}{x_i} p_i^{x_i} (1 - p_i)^{N - x_i}. \quad (2.24)
 \end{aligned}$$

This shows that S_i is just a term of the binomial distribution. By inserting the final expression of Eq. (2.24) in the final expression of (2.23), one gets:

$$\begin{aligned}
 \langle \text{AIVD}(\{p_1, \dots, p_q\}) \rangle &= 1 - \frac{1}{N(N-1)} \sum_{i=1}^q \sum_{x_i \leq N} (x_i^2 - x_i) \binom{N}{x_i} p_i^{x_i} (1 - p_i)^{N - x_i} = \\
 &= 1 - \frac{1}{N(N-1)} \sum_{i=1}^q [N p_i (N p_i - p_i + 1) - N p_i] = 1 - \sum_{i=1}^q p_i^2, \quad (2.25)
 \end{aligned}$$

which concludes the proof of Eq. (2.8), after using the well known expressions for the first and second moments $\langle x_i \rangle$ and $\langle x_i^2 \rangle$ of the binomial distribution. Note that the dependence on N cancels out during the derivation.

Another, arguably shorter proof can be formulated with the aid of indicator functions of the type $\mathbb{I}_i(x)$, which gives 1 if cultural vector x is an entry of trait i and gives 0 otherwise. One can express the feature-level AIVD of one, generic set of cultural vectors in terms of indicator functions and write the expected, feature-level AIVD as an average of this expression. The p_i^2 part of Eq. (2.8) then appears from an averaging of the $\mathbb{I}_i(x)\mathbb{I}_i(y)$ product, where x and y are two arbitrary cultural vectors.

2.C Fitting algorithm

This section explains the procedure used for simultaneously tuning the α and β parameters of either of the two stochastic models of culture, such that a match is obtained between the model and the empirical data, in terms of the averages of the AIVD and SIVD observables:

$$\begin{aligned}
 \langle \text{AIVD}(\alpha, \beta) \rangle &= \text{AIVD}_{\text{emp}}, \\
 \langle \text{SIVD}(\alpha, \beta) \rangle &= \text{SIVD}_{\text{emp}},
 \end{aligned} \quad (2.26)$$

for a fixed number of prototypes k , assuming that either of the two equalities above is satisfied when there is an overlap between the uncertainty range associated to the quantity on the left side and that associated to the quantity on the right side.

There are multiple reasons why this problem is challenging:

- an analytical formula for the $\langle \text{SIVD}(\alpha, \beta) \rangle$ quantity could not be found
- although an analytical formula for the $\langle \text{AIVD}(\alpha, \beta) \rangle$ quantity was found (Eqs. (2.19) to (2.22))⁴, it does not allow for inverting the function and for analytically solving the system
- the $\langle \text{SIVD}(\alpha, \beta) \rangle$, AIVD_{emp} and SIVD_{emp} quantities have non-vanishing uncertainty ranges attached to them

Assuming that there exists a unique solution to the above system, a numerical approach for solving it is in order. The method used here relies on a nested, 2-level, adapted bisection method. The first (inner) level of the method takes care of fitting, via bisection, the first quantity for a fixed β – it finds the α value for which $\langle \text{AIVD}(\alpha, \beta) \rangle = \text{AIVD}_{\text{emp}}$ is satisfied for a given β . The second (outer) level of the method takes care of fitting, via bisection, the second quantity – it finds the β for which $\langle \text{SIVD}(\alpha(\beta), \beta) \rangle = \text{SIVD}_{\text{emp}}$ is satisfied, where $\alpha(\beta)$ is provided by the first level. This choice of assigning the AIVD and SIVD observables and the α and β parameters to the two levels in this manner is numerically convenient for several reasons. First, the AIVD can be much more easily computed via the analytical formula, such that assigning it to the first level, which is repeated multiple times (once for each value of β that the second level samples) is more effective. Second, the model AIVD turns out to be relatively insensitive to β for relatively many combinations of values for the k and α parameters, such that fitting AIVD in terms of α within the first level makes more sense.

In addition to adaptations required by the 2-level scheme, other adaptations with respect to the traditional bisection method are needed for allowing it to work with model and empirical uncertainties, as well as to enhance the numerical precision for the $\langle \text{SIVD}(\alpha, \beta) \rangle$ quantity when needed, to the extent needed. Moreover, in addition to statistical errors originating directly in the empirical uncertainties of the AIVD_{emp} and SIVD_{emp} quantities and in the numerical uncertainty of the model SIVD quantity, the second level of the method is also affected by “systematic errors” on $\langle \text{SIVD}(\alpha(\beta), \beta) \rangle$, originating in the fitting procedure at the first level, and indirectly in the empirical uncertainty of AIVD_{emp} – which for all practical purposes can be assumed fixed, thus motivating using the term “systematic” for its propagation to the model SIVD at the second level.

In order to address all these challenges in a self consistent way, the method developed here turns out to be quite sophisticated, which is why it is explained in detail in the following four sections. Specifically, Sec. 2.C.1 focuses on the first fitting level, Sec. 2.C.2 focuses on the second fitting level, Sec. 2.C.3 describes how

⁴Which implies that the specific uncertainty range of $\langle \text{AIVD}(\alpha, \beta) \rangle$ has a null width.

various sub-problems invoked by the previous two sections are addressed, while Sec. 2.C.4 describes how the tools presented in Sections 2.C.1, 2.C.2 and 2.C.3 are used for producing some of the results presented in Sections 2.3 and 2.4. The method is potentially of use for addressing other problems that are formally similar to the problem presented here, although certain adaptations might be needed.

Since the method has mostly an algorithmic nature, much of it is explained via pseudocode, such that a few conventions that will be extensively used below and that are not necessarily standard are worth mentioning. First, the “=” symbol is used with double meaning: in a normal statement (such as “ $a = b$ ”) it is to be interpreted as an assignment (of the value of variable b to variable a); in the header of an **if** or **while** statement (such as “**if** $a = b$ ”) it is to be interpreted as a check (of whether the values of a and b are equal). A variable is implicitly declared when it first appears, either on the left side of an assignment or in the header of a function definition (in which case it is also called an argument or function parameter); the scope of the variable is the part of the function below and to the right of the place where it first appears. Functions are distinguished from each other through their names, their numbers of arguments and the types of those arguments⁵ On the other hand, the arguments of a function are distinguished from each other via their order. Some variables are actually ordered sequences of other variables, which in turn are denoted by $(x_1, \dots, x_n)$ notation. In the same spirit, an assignments of the type $X = (x_1, \dots, x_n)$ is referred to as a “variable compression”, while one of the type $(x_1, \dots, x_n) = X$ is referred to as a “variable decompression”. These allow for keeping the pseudocode compact, while still rigorous. An uncertainty range refers to an interval $[x - \delta x, x + \delta x]$, where x is a mean and δx is an error relying (directly, or indirectly) on a standard mean error calculation, the uncertainty range being formally encoded by the sorted $(x, \delta x)$ sequence. Note that the square brackets “[.]” are consistently used to denote an interval of real numbers, while the round brackets “(,)” are used to denote an ordered sequence of two or more elements. Finally, it is worth noting that the pseudocode relies heavily on function calls and on recursive definitions, and that there is a certain parallelism between the functions defined in Sec. 2.C.1 and those defined in Sec. 2.C.2.

2.C.1 First level fitting

This section presents the algorithm part concerned with the first fitting level. The algorithm is split in three main functions: FIT-1, BISECT-1, DISPLACE-1, all of them returning the same type of information. FIT-1 always calls BISECT-1, while the latter may or may not call DISPLACE-1 at any stage, which in turn may

⁵Sometimes this can be confusing, since the types of the arguments are only mentioned in the text before the definition of the function. In these cases however, the reader is guided by the names of the arguments, which in the function definition are kept as close as possible to those in the function call(s).

or may not call BISECT-1. The pseudocode also invokes two constants, which are assumed to be known a-priori and available for use anywhere in these three functions. The first constant is $\delta\alpha$, which controls the desired resolution ($\delta\alpha$ is essentially a grid-spacing) in the α parameter, which is here set to the inverse of the number of features: $\delta\alpha = \frac{1}{F}$ ⁶. The second constant is AIVD_{emp} , which stands for the AIVD uncertainty range for the empirical data.

Function FIT-1 acts as an interface for the first-level fitting, which consists of tuning the α parameter, for given values of β and k , such that the AIVD quantity matches the empirical value. Here, β is a real number belonging to $[0, 1]$ while k is a strictly positive integer number. The method returns the left (α_L) and right (α_R) margins of the tightest α interval found, together with the estimated α match within this interval assuming linearity (α_{fit}) and an associated error (α_{err}). It assumes that the empirical AIVD can actually be uniquely matched by varying α , for the given values of β and k . The method essentially carries out some initializations (Lines 2,3), before passing the task to BISECT-1.

```

1: function FIT-1( $\beta, k$ )
2:   ( $\alpha_L, \alpha_R$ ) = Init-1( $\delta\alpha$ )                                ▷ initializing the  $\alpha$ -interval
3:    $\text{AIVD}_L = \langle \text{AIVD} \rangle_{\alpha_L, \beta}^k$ ;  $\text{AIVD}_R = \langle \text{AIVD} \rangle_{\alpha_R, \beta}^k$     ▷ analytics (Eq. (2.19))
4:   return BISECT-1( $\alpha_L, \alpha_R, \text{AIVD}_L, \text{AIVD}_R, \beta, k$ )
5: end function

```

Function BISECT-1 is mostly a typical, recursive implementation of the bisection method. This sequentially narrows down the $[\alpha_L, \alpha_R]$ interval, such that at each stage the empirical AIVD is contained, namely that $\min(\text{AIVD}_L, \text{AIVD}_R) < \text{AIVD}_{\text{emp}} < \max(\text{AIVD}_L, \text{AIVD}_R)$ is satisfied, where the AIVD_L and AIVD_R values correspond to the left and right margins of the α interval. Here, α_L , α_R , AIVD_L and AIVD_R are real numbers belonging to $[0, 1]$ while β and k are of the same type as in FIT-1. It returns the same type of information as FIT-1. The method converges, the fitting being considered complete, when the interval has reached the $\delta\alpha$ resolution limit, in which case estimations for an “ideal” α inside this interval α_{fit} and its error α_{err} are made and returned together with the boundaries of the interval (lines 3-6). Moreover, the method may also call DISPLACE-1 in case the AIVD_M value corresponding to the computed midpoint α_M happens to fall within the AIVD_{emp} uncertainty range (lines 8-10) – this is needed in order to keep the output format consistent and the final α interval relatively narrow. Otherwise, the method decides to zoom in (by calling itself) on either the left or right halves of the interval, depending on the position of AIVD_{emp} with respect to AIVD_L , AIVD_M and AIVD_R (lines 11-16).

```

1: function BISECT-1( $\alpha_L, \alpha_R, \text{AIVD}_L, \text{AIVD}_R, \beta, k$ )
2:    $\alpha_M = \text{MIDDLE}(\alpha_L, \alpha_R, \delta\alpha)$                         ▷ computing midpoint on the  $\alpha$  grid

```

⁶ There is no clear lower bound on $\delta\alpha$, regardless of which stochastic model is used, but $\frac{1}{F}$ is a lower bound on $\delta\beta$ when PG is used, so for simplicity the choice $\delta\alpha = \delta\beta = \frac{1}{F}$ is made.

```

3:  if  $\neg \text{DISTINCT}(\alpha_M, \alpha_L, \alpha_R)$  then
4:       $(\alpha_{\text{fit}}, \alpha_{\text{err}}) = \text{INTERNFITLIN-1}(\alpha_L, \alpha_R, \text{AIVD}_L, \text{AIVD}_R, \text{AIVD}_{\text{emp}})$ 
5:      return  $(\alpha_L, \alpha_R, \alpha_{\text{fit}}, \alpha_{\text{err}})$   $\triangleright$  fitting complete
6:  end if
7:   $\text{AIVD}_M = \langle \text{AIVD} \rangle_{\alpha_M, \beta}^k$   $\triangleright$  analytics (Eq. (2.19))
8:  if  $\text{MATCH-1}(\text{AIVD}_M, \text{AIVD}_{\text{emp}})$  then
9:      return  $\text{DISPLACE-1}(\alpha_L, \alpha_R, \alpha_M, \text{AIVD}_L, \text{AIVD}_R, \beta, k)$ 
10: end if
11: if  $\text{ORD-1}(\text{AIVD}_L, \text{AIVD}_R) = \text{ORD-1}(\text{AIVD}_M, \text{AIVD}_{\text{emp}})$  then
12:      $\alpha_L = \alpha_M; \text{AIVD}_L = \text{AIVD}_M$   $\triangleright$  selecting right interval
13: else
14:      $\alpha_R = \alpha_M; \text{AIVD}_R = \text{AIVD}_M$   $\triangleright$  selecting left interval
15: end if
16: return  $\text{BISECT-1}(\alpha_L, \alpha_R, \text{AIVD}_L, \text{AIVD}_R, \beta, k)$   $\triangleright$  zooming in
17: end function
    
```

Function `DISPLACE-1` attempts to displace the midpoint α_M previously calculated at some stage in `BISECT-1`, in a way that its associated `AIVD` would fall outside the empirical uncertainty range. This function has all the arguments of `BISECT-1` and α_M as an additional one, which is a real number belonging to $[0, 1]$. It returns the same type of information as `FIT-1`. The method first computes a “secondary” midpoint α'_M to the left of α_M and its corresponding AIVD'_M value. If the resolution limit $\delta\alpha$ is not reached and AIVD'_M falls outside the AIVD_{emp} range, `BISECT-1` is applied further to the $[\alpha'_M, \alpha_R]$ interval (lines 2-11). Otherwise, the analogous procedure is applied on the right side (12-21). If the procedure fails to provide a convenient, secondary midpoint on either side, the fitting is considered complete with the current $[\alpha_L, \alpha_R]$ interval and the $\alpha_{\text{fit}}, \alpha_{\text{err}}$ estimates made like in `BISECT-1` (lines 22-23).

```

1: function DISPLACE-1( $\alpha_L, \alpha_R, \alpha_M, \text{AIVD}_L, \text{AIVD}_R, \beta, k$ )
2:    $\alpha'_M = \text{MIDDLE}(\alpha_L, \alpha_M, \delta\alpha)$   $\triangleright$  trying displacement to the left
3:   if  $\text{DISTINCT}(\alpha'_M, \alpha_L, \alpha_M)$  then
4:        $\text{AIVD}'_M = \langle \text{AIVD} \rangle_{\alpha'_M, \beta}^k$   $\triangleright$  analytics (Eq. (2.19))
5:       if  $\neg \text{MATCH-1}(\text{AIVD}'_M, \text{AIVD}_{\text{emp}})$  then
6:           if  $\text{ORD-1}(\text{AIVD}_L, \text{AIVD}_R) = \text{ORD-1}(\text{AIVD}'_M, \text{AIVD}_{\text{emp}})$  then
7:                $\alpha_L = \alpha'_M; \text{AIVD}_L = \text{AIVD}'_M$ 
8:               return  $\text{BISECT-1}(\alpha_L, \alpha_R, \text{AIVD}_L, \text{AIVD}_R, \beta, k)$   $\triangleright$  zooming in
9:           end if
10:       end if
11:   end if
12:    $\alpha'_M = \text{MIDDLE}(\alpha_M, \alpha_R, \delta\alpha)$   $\triangleright$  trying displacement to the right
13:   if  $\text{DISTINCT}(\alpha'_M, \alpha_M, \alpha_R)$  then
14:        $\text{AIVD}'_M = \langle \text{AIVD} \rangle_{\alpha'_M, \beta}^k$   $\triangleright$  analytics (Eq. (2.19))
15:       if  $\neg \text{MATCH-1}(\text{AIVD}'_M, \text{AIVD}_{\text{emp}})$  then
    
```

```

16:         if ORD-1(AIVDL, AIVDR) ≠ ORD-1(AIVD'M, AIVDemp) then
17:             αR = α'M; AIVDR = AIVD'M
18:             return BISECT-1(αL, αR, AIVDL, AIVDR, β, k) ▷ zooming in
19:         end if
20:     end if
21: end if
22: (αfit, αerr) = INTERNFITLIN-1(αL, αR, AIVDL, AIVDR, AIVDemp)
23: return (αL, αR, αfit, αerr) ▷ fitting complete
24: end function

```

2.C.2 Second level fitting

This section presents the algorithm part concerned with the second fitting level. Each of the three functions of the first fitting level (Sec. 2.C.1) has a correspondent here: FIT-2, BISECT-2, DISPLACE-2, all of them returning the same type of information⁷, each of them having a similar, structure, purpose and role to the correspondent within the first fitting level. Additionally, this section presents the pseudocode for a fourth function, NUMSIVD, which carries out the numerical SIVD calculations. In addition to the two constants introduced at the first level, the second level pseudocode invokes two other constants, which are also assumed to be known a-priori and available for use anywhere in these four functions. First, $\delta\beta$ is the desired resolution in the β parameter, which is here set to the inverse of the number of features: $\delta\beta = \frac{1}{F}$. Second, SIVD_{emp} is the SIVD uncertainty range for the empirical data.

In relation to the first three functions, the descriptions below attempt to mostly emphasize the elements that come in addition with respect to their first-level correspondents. Some of these elements have a repetitive nature and are worth explaining before moving to the specific description of each function. First, the (generic) $\tilde{\beta}_X$ notation (where “X” can stand for “L”, “R” or “M”) denotes the (generic) “composite fitting information” $\tilde{\beta}_X = (\beta, \alpha_L, \alpha_R, \alpha_{\text{fit}}, \alpha_{\text{err}})_X$, which is a 5-tuple consisting of a β value together with the associated four values returned by a (generic) call FIT-1(β, k) for that specific β and some arbitrary k . Second, whenever an “SIVD_X” variable appears in the first three functions (where “X” is again a generic label), except for SIVD_{emp}, it actually denotes the (generic) “composite SIVD information” $\text{SIVD}_X = ((\text{SIVD}_L^{\text{fit}}, \text{SIVD}_L^{\text{err}}), (\text{SIVD}_R^{\text{fit}}, \text{SIVD}_R^{\text{err}}))_X$, which is a pair of pairs of real numbers, each inner pair corresponding to a model SIVD uncertainty range associated to one margin of an α interval returned by a call to FIT-1, while both inner pairs have the same β . This schematically reads:

$$\begin{aligned}
(\beta, \alpha_L) &\rightarrow (\text{SIVD}_L^{\text{fit}}, \text{SIVD}_L^{\text{err}}), \\
(\beta, \alpha_R) &\rightarrow (\text{SIVD}_R^{\text{fit}}, \text{SIVD}_R^{\text{err}}),
\end{aligned}$$

⁷The type of information returned by the three functions at a second-level fitting is different than that of the three functions at the first-level fitting, and actually more complex.

Third, any (generic) call $\text{NUMSIVD}(\bar{\beta}, k)$ is necessarily preceded by an associated (generic) call $\text{FIT-1}(\beta, k)$ and by an associated (generic) variable compression $\bar{\beta} = (\beta, \alpha_L, \alpha_R, \alpha_{\text{fit}}, \alpha_{\text{err}})$, the last two being needed for producing the composite fitting information $\bar{\beta}$. Fourth, whenever a piece of composite SIVD information appears in a call to ORD-2 or MATCH-2 , it is accompanied by an associated piece of composite fitting information, which allows for the mean, statistical error and systematic error of in the model SIVD to be all reconstructed within, for a given combination of β and k .

Function FIT-2 acts as an interface for the second-level fitting, which consists of tuning the β parameter, for a given value of k , such that the SIVD quantity matches the empirical value, relying on an underlying tuning of the α parameter in terms of the AIVD quantity (using FIT-1). Here, k is a strictly positive, integer number. The method returns the composite fitting information associated to the left ($\bar{\beta}_L$) and right ($\bar{\beta}_R$) margins of the tightest β interval found, together with the estimated β match within this interval (β_{fit}) and its associated error (β_{err}). It assumes that the empirical SIVD can actually be uniquely matched by varying β and α , for the given value of k . After checking that there exists a meaningful $[\beta_L, \beta_R]$ interval for which the first-level fitting is possible (lines 2,3), the method conducts the numeric SIVD calculations on both sides of the interval (line 6), preceded, on each side, by the first level fitting and the decompression (lines 4,5, as explained above), in order to finally pass the task to BISECT-2 .

```

1: function FIT-2( $k$ )
2:    $(\beta_L, \beta_R) = \text{Init-2}(\delta\beta, k, \text{AIVD}_{\text{emp}})$  ▷ initializing the  $\beta$ -interval
3:   if  $\beta_L < \beta_R$  then
4:      $(\alpha_L^L, \alpha_L^R, \alpha_L^{\text{fit}}, \alpha_L^{\text{err}}) = \text{FIT-1}(\beta_L, k); (\alpha_R^L, \alpha_R^R, \alpha_R^{\text{fit}}, \alpha_R^{\text{err}}) = \text{FIT-1}(\beta_R, k)$ 
5:      $\bar{\beta}_L = (\beta_L, \alpha_L^L, \alpha_L^R, \alpha_L^{\text{fit}}, \alpha_L^{\text{err}}); \bar{\beta}_R = (\beta_R, \alpha_R^L, \alpha_R^R, \alpha_R^{\text{fit}}, \alpha_R^{\text{err}})$ 
6:      $\text{SIVD}_L = \text{NUMSIVD}(\bar{\beta}_L, k); \text{SIVD}_R = \text{NUMSIVD}(\bar{\beta}_R, k)$  ▷ numerics
7:     return  $\text{BISECT-2}(\bar{\beta}_L, \bar{\beta}_R, \text{SIVD}_L, \text{SIVD}_R, k)$ 
8:   end if
9:   return  $\text{FittingImpossibleError}$ 
10: end function
    
```

Function BISECT-2 is another recursive implementation of the bisection method, which sequentially narrows down the $[\beta_L, \beta_R]$ interval, such that at each stage the empirical SIVD is contained. Here, $\bar{\beta}_L, \bar{\beta}_R$ are 5-tuples of real numbers encoding the left and right pieces of composite fitting information, $\text{SIVD}_L, \text{SIVD}_R$ are the pairs of pairs of real numbers encoding the left- β and right- β pieces of composite SIVD information, while k is of the same type as in FIT-2 . It returns the same type of information as FIT-2 . Like BISECT-1 , the function consists of a part concerned with convergence (lines 4-7), a part concerned with the jump to DISPLACE-2 (lines 11-13) and a part concerned with choosing between the left and right β subintervals and with zooming in on the chosen one (lines 14-19). Note the additional statements concerned with decompressing the composite fit-

ting information (line 2) and with preparing the numeric SIVD calculations at the midpoint (lines 8-9).

```

1: function BISECT-2( $\bar{\beta}_L, \bar{\beta}_R, \text{SIVD}_L, \text{SIVD}_R, k$ )
2:    $(\beta_L, \alpha_L^L, \alpha_L^R, \alpha_L^{\text{fit}}, \alpha_L^{\text{err}}) = \bar{\beta}_L$ ;  $(\beta_R, \alpha_R^L, \alpha_R^R, \alpha_R^{\text{fit}}, \alpha_R^{\text{err}}) = \bar{\beta}_R$ 
3:    $\beta_M = \text{MIDDLE}(\beta_L, \beta_R, \delta\beta)$ 
4:   if  $\neg \text{DISTINCT}(\beta_M, \beta_L, \beta_R)$  then
5:      $(\beta_{\text{fit}}, \beta_{\text{err}}) = \text{INTERNFITLIN-2}(\bar{\beta}_L, \bar{\beta}_R, \text{SIVD}_L, \text{SIVD}_R, \text{SIVD}_{\text{emp}})$ 
6:     return  $(\bar{\beta}_L, \bar{\beta}_R, \beta_{\text{fit}}, \beta_{\text{err}})$ 
7:   end if
8:    $(\alpha_M^L, \alpha_M^R, \alpha_M^{\text{fit}}, \alpha_M^{\text{err}}) = \text{FIT-1}(\beta_M, k)$ 
9:    $\bar{\beta}_M = (\beta_M, \alpha_M^L, \alpha_M^R, \alpha_M^{\text{fit}}, \alpha_M^{\text{err}})$ 
10:   $\text{SIVD}_M = \text{NUMSIVD}(\bar{\beta}_M, k)$  ▷ numerics
11:  if  $\text{MATCH-2}(\bar{\beta}_M, \text{SIVD}_M, \text{SIVD}_{\text{emp}})$  then
12:    return  $\text{DISPLACE-2}(\bar{\beta}_L, \bar{\beta}_R, \bar{\beta}_M, \text{SIVD}_L, \text{SIVD}_R, k)$ 
13:  end if
14:  if  $\text{ORD-2}(\bar{\beta}_L, \bar{\beta}_R, \text{SIVD}_L, \text{SIVD}_R) = \text{ORD-2}(\bar{\beta}_M, \text{SIVD}_M, \text{SIVD}_{\text{emp}})$  then
15:     $\bar{\beta}_L = \bar{\beta}_M$ ;  $\text{SIVD}_L = \text{SIVD}_M$  ▷ selecting right interval
16:  else
17:     $\bar{\beta}_R = \bar{\beta}_M$ ;  $\text{SIVD}_R = \text{SIVD}_M$  ▷ selecting left interval
18:  end if
19:  return  $\text{BISECT-2}(\bar{\beta}_L, \bar{\beta}_R, \text{SIVD}_L, \text{SIVD}_R, k)$  ▷ zooming in
20: end function

```

Function `DISPLACE-2` attempts to displace the midpoint β_M previously calculated at some stage in `BISECT-2`, in a way that its associated SIVD uncertainty range does not overlap with the empirical one. This function has all the arguments of `BISECT-1` and $\bar{\beta}_M$ as an additional one, which is a 5-tuple of real numbers encoding the midpoint composite fitting information. It returns the same type of information as `FIT-2`. Like `DISPLACE-1`, the function consists of a part that attempts a displacement to the left (lines 4-15), one that attempts a displacement to the right (lines 16-27) and one that takes care of the convergence (lines 28-29). Note the additional statements concerned with decompressing the composite fitting information (lines 2-3) and with preparing the numeric SIVD calculations for the left/right secondary midpoint (lines 6-7/18-19).

```

1: function DISPLACE-2( $\bar{\beta}_L, \bar{\beta}_R, \bar{\beta}_M, \text{SIVD}_L, \text{SIVD}_R, k$ )
2:    $(\beta_L, \alpha_L^L, \alpha_L^R, \alpha_L^{\text{fit}}, \alpha_L^{\text{err}}) = \bar{\beta}_L$ ;  $(\beta_R, \alpha_R^L, \alpha_R^R, \alpha_R^{\text{fit}}, \alpha_R^{\text{err}}) = \bar{\beta}_R$ ;
3:    $(\beta_M, \alpha_M^L, \alpha_M^R, \alpha_M^{\text{fit}}, \alpha_M^{\text{err}}) = \bar{\beta}_M$ 
4:    $\beta'_M = \text{MIDDLE}(\beta_L, \beta_M, \delta\beta)$  ▷ trying displacement to the left
5:   if  $\text{DISTINCT}(\beta'_M, \beta_L, \beta_M)$  then
6:      $(\dot{\alpha}_M^L, \dot{\alpha}_M^R, \dot{\alpha}_M^{\text{fit}}, \dot{\alpha}_M^{\text{err}}) = \text{FIT-1}(\beta'_M, k)$ 
7:      $\bar{\beta}'_M = (\beta'_M, \dot{\alpha}_M^L, \dot{\alpha}_M^R, \dot{\alpha}_M^{\text{fit}}, \dot{\alpha}_M^{\text{err}})$ 
8:      $\text{SIVD}'_M = \text{NUMSIVD}(\bar{\beta}'_M, k)$  ▷ numerics

```

```

9:      if  $\neg \text{MATCH-2}(\bar{\beta}'_M, \text{SIVD}'_M, \text{SIVD}_{\text{emp}})$  then
10:         if  $\text{ORD-2}(\bar{\beta}_L, \bar{\beta}_R, \text{SIVD}_L, \text{SIVD}_R) = \text{ORD-2}(\bar{\beta}'_M, \text{SIVD}'_M, \text{SIVD}_{\text{emp}})$ 
then
11:             $\bar{\beta}_L = \bar{\beta}'_M$ ;  $\text{SIVD}_L = \text{SIVD}'_M$ 
12:            return  $\text{BISECT-2}(\bar{\beta}_L, \bar{\beta}_R, \text{SIVD}_L, \text{SIVD}_R, k)$   $\triangleright$  zooming in
13:         end if
14:      end if
15:   end if
16:    $\beta'_M = \text{MIDDLE}(\beta_M, \beta_R, \delta\beta)$   $\triangleright$  trying displacement to the right
17:   if  $\text{DISTINCT}(\beta'_M, \beta_M, \beta_R)$  then
18:       $(\dot{\alpha}_M^L, \dot{\alpha}_M^R, \dot{\alpha}_M^{\text{fit}}, \dot{\alpha}_M^{\text{err}}) = \text{FIT-1}(\beta'_M, k)$ 
19:       $\bar{\beta}'_M = (\beta'_M, \dot{\alpha}_M^L, \dot{\alpha}_M^R, \dot{\alpha}_M^{\text{fit}}, \dot{\alpha}_M^{\text{err}})$ 
20:       $\text{SIVD}'_M = \text{NUMSIVD}(\bar{\beta}'_M, k)$   $\triangleright$  numerics
21:      if  $\neg \text{MATCH-2}(\bar{\beta}'_M, \text{SIVD}'_M, \text{SIVD}_{\text{emp}})$  then
22:         if  $\text{ORD-2}(\bar{\beta}_L, \bar{\beta}_R, \text{SIVD}_L, \text{SIVD}_R) \neq \text{ORD-2}(\bar{\beta}'_M, \text{SIVD}'_M, \text{SIVD}_{\text{emp}})$ 
then
23:             $\bar{\beta}_R = \bar{\beta}'_M$ ;  $\text{SIVD}_R = \text{SIVD}'_M$ 
24:            return  $\text{BISECT-2}(\bar{\beta}_L, \bar{\beta}_R, \text{SIVD}_L, \text{SIVD}_R, k)$   $\triangleright$  zooming in
25:         end if
26:      end if
27:   end if
28:    $(\beta_{\text{fit}}, \beta_{\text{err}}) = \text{INTERNFITLIN-2}(\bar{\beta}_L, \bar{\beta}_R, \text{SIVD}_L, \text{SIVD}_R, \text{SIVD}_{\text{emp}})$ 
29:   return  $(\bar{\beta}_L, \bar{\beta}_R, \beta_{\text{fit}}, \beta_{\text{err}})$ 
30: end function

```

Function NUMSIVD numerically generates a piece of composite SIVD information with a precision that is as high as possible. Here, $\bar{\beta}$ is a 5-tuple of real numbers encoding a composite fitting information, while k is a positive integer number. One sequence of SIVD values is numerically generated (lines 4-5 and 12-13) for each of the two margins of the α interval (contained in $\bar{\beta}$), for the given β (also contained in $\bar{\beta}$) and the given k . An uncertainty range is obtained from each of the two sequences (lines 6 and 16). These two uncertainty ranges are used together with the information in $\bar{\beta}$ to produce estimates for an average, a statistical error and a systematic error that are β -specific rather than (α, β) -specific (lines 7,8 and 17,18). The number of SIVD values in the two sequences is increased and the calculations are repeated as long as the condition in line 10 remains true, namely as long as: the statistical error is higher than the systematic error, the desired separation between the model and empirical (statistical) uncertainty ranges is not reached and the maximal SIVD sequence length is not reached. The desired separation and the SIVD sequence length are controlled via variables s and n , initialized in line 2 – the initial values of these variables, as well as the upper bound on the latter are hard-coded, as visible in the pseudocode, and have been decided after some experimentation with NUMSIVD, but they are not essential for the actual outcome. Also note the decompression of the composite

fitting information (line 3) and the decompression of SIVD uncertainty ranges (lines 9 and 19).

```

1: function NUMSIVD( $\bar{\beta}, k$ )
2:    $n = 20; s = 5$   $\triangleright$  initial number of realizations and desired separation
3:    $(\beta, \alpha_L, \alpha_R, \alpha_{\text{fit}}, \alpha_{\text{err}}) = \bar{\beta}$ 
4:    $\text{SIVD}_L^{\text{seq}} = \text{GENSEQSIVD}(\alpha_L, \beta, k, n)$ 
5:    $\text{SIVD}_R^{\text{seq}} = \text{GENSEQSIVD}(\alpha_R, \beta, k, n)$ 
6:    $\text{SIVD}_L = \text{COMPAVGERR}(\text{SIVD}_L^{\text{seq}}); \text{SIVD}_R = \text{COMPAVGERR}(\text{SIVD}_R^{\text{seq}})$ 
7:    $\text{SIVD} = \text{INTERPOL}(\alpha_L, \alpha_R, \alpha_{\text{fit}}, \text{SIVD}_L, \text{SIVD}_R)$ 
8:    $\text{SIVD}_{\text{syst}} = \text{COMPSYSTERR}(\alpha_L, \alpha_R, \alpha_{\text{err}}, \text{SIVD}_L, \text{SIVD}_R)$ 
9:    $(\text{SIVD}_{\text{avg}}, \text{SIVD}_{\text{stat}}) = \text{SIVD}; (\text{SIVD}_{\text{emp}}^{\text{avg}}, \text{SIVD}_{\text{emp}}^{\text{stat}}) = \text{SIVD}_{\text{emp}}$ 
10:  while  $\text{SIVD}_{\text{stat}} > \text{SIVD}_{\text{syst}} \wedge (\text{SIVD}_{\text{stat}} + \text{SIVD}_{\text{emp}}^{\text{stat}} > |\text{SIVD}_{\text{emp}}^{\text{avg}} - \text{SIVD}_{\text{avg}}|/s) \wedge$ 
     $n < 350$  do
11:     $n = 2 \cdot n$ 
12:     $\text{SIVD}_L^{\text{tmpSeq}} = \text{GENSEQSIVD}(\alpha_L, \beta, k, n)$ 
13:     $\text{SIVD}_R^{\text{tmpSeq}} = \text{GENSEQSIVD}(\alpha_R, \beta, k, n)$ 
14:     $\text{SIVD}_L^{\text{seq}} = \text{MERGE}(\text{SIVD}_L^{\text{seq}}, \text{SIVD}_L^{\text{tmpSeq}})$ 
15:     $\text{SIVD}_R^{\text{seq}} = \text{MERGE}(\text{SIVD}_R^{\text{seq}}, \text{SIVD}_R^{\text{tmpSeq}})$ 
16:     $\text{SIVD}_L = \text{COMPAVGERR}(\text{SIVD}_L^{\text{seq}}); \text{SIVD}_R = \text{COMPAVGERR}(\text{SIVD}_R^{\text{seq}})$ 
17:     $\text{SIVD} = \text{INTERPOL}(\alpha_L, \alpha_R, \alpha_{\text{fit}}, \text{SIVD}_L, \text{SIVD}_R)$ 
18:     $\text{SIVD}_{\text{syst}} = \text{COMPSYSTERR}(\alpha_L, \alpha_R, \alpha_{\text{err}}, \text{SIVD}_L, \text{SIVD}_R)$ 
19:     $(\text{SIVD}_{\text{avg}}, \text{SIVD}_{\text{stat}}) = \text{SIVD}$ 
20:  end while
21:  return  $(\text{SIVD}_L, \text{SIVD}_R)$ 
22: end function

```

2.C.3 Used functions

This section describes functions that are used by the pseudocode in sections 2.C.1 or 2.C.2 but are not described there. The following is a list of functions for which the pseudocode is also provided, following each text description.

Function INTERFITLIN-1 fine-tunes the α parameter such that AIVD_{emp} is matched, relying on a linear approximation of the model AIVD as a function of α within the (α_L, α_R) interval, using the boundary values AIVD_L and AIVD_R . Its arguments are of the same type as those of INTERFITLIN (described below), except that AIVD_L and AIVD_R are real numbers rather than uncertainty ranges. The output structure is entirely the same as that of INTERFITLIN. It is essentially a first-level fitting interface for INTERNFITLIN, which is called after specifying that the errors associated to AIVD_L and AIVD_R are zero.

```

1: function INTERNFITLIN-1( $\alpha_L, \alpha_R, \text{AIVD}_L, \text{AIVD}_R, \text{AIVD}_{\text{emp}}$ )
2:    $\text{AIVD}'_L = (\text{AIVD}_L, 0); \text{AIVD}'_R = (\text{AIVD}_R, 0)$ 

```

```

3:   return INTERFITLIN( $\alpha_L, \alpha_R, \text{AIVD}'_L, \text{AIVD}'_R, \text{AIVD}_{\text{emp}}$ )
4: end function

```

Function INTERFITLIN-2 fine-tunes the β parameter such that SIVD_{emp} is matched, relying on a linear approximation of the model SIVD as a function of β within the $[\beta_L, \beta_R]$ interval, using the boundary information stored in SIVD_L and SIVD_R . Its arguments are of the same type as those of INTERFITLIN (described below), except that $\bar{\beta}_L$ and $\bar{\beta}_R$ are 5-tuples or real numbers rather than real numbers and SIVD_L and SIVD_R are pieces composite SIVD information rather than uncertainty ranges. The output structure is entirely the same as that of INTERFITLIN. It is essentially a second-level fitting interface for INTERFITLIN, which is called after carrying out the following two operations: computing the mean, statistical error and systematic error on each of the two margins of the β interval, using the right combination of composite fitting information and composite SIVD information (lines 2,3); compressing information into an SIVD uncertainty range for each of the two margins, after choosing the highest among the two errors for each margin.

```

1: function INTERFITLIN-2( $\bar{\beta}_L, \bar{\beta}_R, \text{SIVD}_L, \text{SIVD}_R, \text{SIVD}_{\text{emp}}$ )
2:   ( $\text{SIVD}_L^{\text{avg}}, \text{SIVD}_L^{\text{stat}}, \text{SIVD}_L^{\text{syst}}$ ) = MEANSTATSYST( $\bar{\beta}_L, \text{SIVD}_L$ )
3:   ( $\text{SIVD}_R^{\text{avg}}, \text{SIVD}_R^{\text{stat}}, \text{SIVD}_R^{\text{syst}}$ ) = MEANSTATSYST( $\bar{\beta}_R, \text{SIVD}_R$ )
4:    $\text{SIVD}'_L = (\text{SIVD}_L^{\text{avg}}, \text{MAX}(\text{SIVD}_L^{\text{stat}}, \text{SIVD}_L^{\text{syst}}))$ 
5:    $\text{SIVD}'_R = (\text{SIVD}_R^{\text{avg}}, \text{MAX}(\text{SIVD}_R^{\text{stat}}, \text{SIVD}_R^{\text{syst}}))$ 
6:   return INTERFITLIN( $\beta_L, \beta_R, \text{SIVD}'_L, \text{SIVD}'_R, \text{SIVD}_{\text{emp}}$ )
7: end function

```

Function MATCH-1 checks whether AIVD (real value) falls within the uncertainty range specified by AIVD_{emp} . It acts as an interface for MATCH (described below) within the first-level fitting scheme.

```

1: function MATCH-1(AIVD,  $\text{AIVD}_{\text{emp}}$ )
2:    $\text{AIVD}' = (\text{AIVD}, 0)$ 
3:   return MATCH( $\text{AIVD}'$ ,  $\text{AIVD}_{\text{emp}}$ )
4: end function

```

Function MATCH-2 checks whether there is an overlap between the model SIVD uncertainty range obtained from $\bar{\beta}$ (composite fitting information) and SIVD (composite SIVD information) and the empirical one encoded by SIVD_{emp} . It acts as an interface for MATCH (described below) within the second-level fitting scheme.

```

1: function MATCH-2( $\bar{\beta}, \text{SIVD}, \text{SIVD}_{\text{emp}}$ )
2:   ( $\text{SIVD}_{\text{avg}}, \text{SIVD}_{\text{stat}}, \text{SIVD}_{\text{syst}}$ ) = MEANSTATSYST( $\bar{\beta}, \text{SIVD}$ )
3:    $\text{SIVD}' = (\text{SIVD}_{\text{avg}}, \text{MAX}(\text{SIVD}_{\text{stat}}, \text{SIVD}_{\text{syst}}))$ 
4:   return MATCH( $\text{SIVD}'$ ,  $\text{SIVD}_{\text{emp}}$ )

```

5: **end function**

Function ORD-1 (first version) checks whether AIVD_L (real value) is smaller than AIVD_R (real value), acting as an interface for ORD within the first-level fitting scheme.

```

1: function ORD-1( $\text{AIVD}_L, \text{AIVD}_R$ )
2:   return ORD( $\text{AIVD}_L, \text{AIVD}_R$ )
3: end function

```

Function ORD-1 (second version) checks whether AIVD (real value) is smaller than the average stored in AIVD_{emp} (uncertainty range), acting as an interface for ORD within the first-level fitting scheme.

```

1: function ORD-1( $\text{AIVD}, \text{AIVD}_{\text{emp}}$ )
2:   ( $\text{AIVD}_{\text{emp}}^{\text{avg}}, \text{AIVD}_{\text{emp}}^{\text{err}}$ ) =  $\text{AIVD}_{\text{emp}}$ 
3:   return ORD( $\text{AIVD}, \text{AIVD}_{\text{emp}}^{\text{avg}}$ )
4: end function

```

Function ORD-2 (first version) checks whether the average stored in the SIVD uncertainty range obtained from $\bar{\beta}_L$ (composite fitting information) and SIVD_L (composite SIVD information) is smaller than the average stored in that obtained from $\bar{\beta}_R$ (composite fitting information) and SIVD_R (composite SIVD information), acting as an interface for ORD within the second-level fitting scheme.

```

1: function ORD-2( $\bar{\beta}_L, \bar{\beta}_R, \text{SIVD}_L, \text{SIVD}_R$ )
2:   ( $\text{SIVD}_L^{\text{avg}}, \text{SIVD}_L^{\text{stat}}, \text{SIVD}_L^{\text{synt}}$ ) = MEANSTATSYST( $\bar{\beta}_L, \text{SIVD}_L$ )
3:   ( $\text{SIVD}_R^{\text{avg}}, \text{SIVD}_R^{\text{stat}}, \text{SIVD}_R^{\text{synt}}$ ) = MEANSTATSYST( $\bar{\beta}_R, \text{SIVD}_R$ )
4:   return ORD( $\text{SIVD}_L^{\text{avg}}, \text{SIVD}_R^{\text{avg}}$ )
5: end function

```

Function ORD-2 (second version) checks whether the average stored in the SIVD uncertainty range obtained from $\bar{\beta}$ (composite fitting information) and SIVD (composite SIVD information) is smaller than the average stored SIVD_{emp} , acting as an interface for ORD within the second-level fitting scheme.

```

1: function ORD-2( $\bar{\beta}, \text{SIVD}, \text{SIVD}_{\text{emp}}$ )
2:   ( $\text{SIVD}_{\text{avg}}, \text{SIVD}_{\text{stat}}, \text{SIVD}_{\text{synt}}$ ) = MEANSTATSYST( $\bar{\beta}, \text{SIVD}$ )
3:   ( $\text{AIVD}_{\text{emp}}^{\text{avg}}, \text{AIVD}_{\text{emp}}^{\text{err}}$ ) =  $\text{AIVD}_{\text{emp}}$ 
4:   return ORD( $\text{SIVD}_{\text{avg}}, \text{AIVD}_{\text{emp}}^{\text{avg}}$ )
5: end function

```

Function MEANSTATSYST estimates a mean, a statistical error and a systematic error from a piece of composite fitting information and an associated piece of composite SIVD information, which are the two arguments of the function. It returns the 3-tuple comprising of the three computed real numbers. Note the decompression of composite fitting information (line 2) and the decompression of

composite SIVD information (line 3).

```

1: function MEANSTATSYST( $\bar{\beta}$ , SIVD)
2:   ( $\beta, \alpha_L, \alpha_R, \alpha_{\text{fit}}, \alpha_{\text{err}}$ ) =  $\bar{\beta}$ 
3:   (SIVDL, SIVDR) = SIVD
4:   SIVD' = INTERPOL( $\alpha_L, \alpha_R, \alpha_{\text{fit}}$ , SIVDL, SIVDR)
5:   SIVDsyst = COMPSYSTERR( $\alpha_L, \alpha_R, \alpha_{\text{err}}$ , SIVDL, SIVDR)
6:   (SIVDavg, SIVDstat) = SIVD'
7:   return (SIVDavg, SIVDstat, SIVDsyst)
8: end function
    
```

The following is a list of functions for which only text explanations are provided in schematic way, sometimes accompanied by figures.

- INIT-1($\delta\alpha$):
 - gives the left and right boundaries of the largest possible interval for which the α parameter is compatible with the stochastic model in use, given the grid spacing $\delta\alpha$
 - input: δ is a real number
 - in practice it returns $(\delta\alpha, 1 - \delta\alpha)$ regardless of whether PG or MPG is used
- INIT-2($\delta\beta, k, \text{AIVD}_{\text{emp}}$):
 - gives the left and right boundaries of the largest possible interval, if any, for which the β parameter allows for the (first level) fitting of $\text{AIVD}(\alpha)$ to successfully take place, given the grid spacing $\delta\beta$
 - input: $\delta\beta$ is a real number, k is a positive integer and AIVD_{emp} is an uncertainty range
 - assumes that there exists at most one β interval $[\beta_L, \beta_R]$ for which there exists an α such that $\langle \text{AIVD} \rangle_{\alpha, \beta}^k = \text{AIVD}_{\text{emp}}$ is satisfied
 - starts from the largest interval allowed by the model and independently adjusts each of the two boundaries via a branching algorithm, until the desired interval is reached
 - returns two (incompatible) boundaries $\beta_L > \beta_R$ if such an interval does not exist
- MIDDLE(l, r, δ):
 - computes the value closest to the average between l and r , on a grid of spacing δ
 - input: l, r, δ are all real numbers
 - assumes that the interval length $l - r$ is equal to an integer times δ

- **DISTINCT(m, l, r):**
 - checks whether m is different than both l and r
 - input: m, l, r are all real numbers constrained to a grid of constant spacing
- **INTERNFITLIN($p_L, p_R, O_L, O_R, O_{\text{emp}}$):**
 - adjusts a parameter p such that an observable O attains a value compatible with the empirical in O_{emp} interval, assuming that O is a linear function of p within the $[p_L, p_R]$ interval
 - input: p_L, p_R are real numbers, encoding the left and right boundaries of the interval; O_L, O_R are mean-error pairs of real numbers encoding the theoretical uncertainty ranges of the observable for the left and for the right boundaries; O_{emp} is a mean-error pair of real numbers encoding the empirical uncertainty range
 - returns the value and associated error of the p parameter resulting from this fitting process ($p_{\text{fit}}, p_{\text{err}}$), computed based on geometrical considerations, in the manner illustrated in Fig. 2.7(a)
 - p_{fit} is calculated first by intersecting the theoretical line with the empirical one, disregarding all errors; then, p_{err} is calculated by assuming that the theoretical error is constant within the $[p_L, p_R]$ interval, with value given by interpolating the errors contained by O_L and O_R at p_{fit}
 - p_{err} takes its origin both in the the empirical error as well as in the theoretical error, but also depends on the slope resulting from the linear approximation
- **MATCH(r_1, r_2):**
 - checks whether there is an overlap between the uncertainty ranges encoded by r_1 and r_2
 - input: r_1, r_2 are mean-error pairs of real numbers
- **ORD(v_L, v_R):**
 - checks whether the condition $v_L < v_R$ is satisfied
 - input: v_L, v_R are real numbers
 - assumes that $v_L \neq v_R$
- **GENSEQSIVD(α, β, k, n)**
 - numerically generates a sequence of n SIVD values according to the respective stochastic model, subject to parameter values indicated by k, α, β

- input: α, β are real numbers, while k, n are positive integers
- $\text{MERGE}(\text{SIVD}_1^{\text{seq}}, \text{SIVD}_2^{\text{seq}})$
 - merges two sequences of (real) SIVD values
 - input: $\text{SIVD}_1^{\text{seq}}, \text{SIVD}_2^{\text{seq}}$ are both sequences of (real) SIVD values
- $\text{COMPAVGERR}(\text{SIVD}^{\text{seq}})$
 - computes the mean and standard error of the mean from SIVD^{seq}
 - input: SIVD^{seq} is a sequence of real SIVD values
- $\text{INTERPOL}(\alpha_L, \alpha_R, \alpha_{\text{fit}}, \text{SIVD}_L, \text{SIVD}_R)$
 - estimates the mean and error in SIVD corresponding to α_{fit} based on the values attained for α_L
 - input: $\alpha_L, \alpha_R, \alpha_{\text{fit}}$ are real numbers, while $\text{SIVD}_L, \text{SIVD}_R$ are mean-error pairs of real numbers
 - uses on a linear interpolation within the $[\alpha_L, \alpha_R]$ interval, separately for the mean and for the error
- $\text{COMPSYSEERR}(\alpha_L, \alpha_R, \alpha_{\text{err}}, \text{SIVD}_L, \text{SIVD}_R)$
 - estimates the systematic error $\text{SIVD}^{\text{syst}}$ of the SIVD quantity induced by the error α_{err} (associated to fitting the α parameter in terms of the AIVD quantity), assuming that SIVD is a linear function of α within the $[\alpha_L, \alpha_R]$ interval
 - input: $\alpha_L, \alpha_R, \alpha_{\text{err}}$ are real numbers while $\text{SIVD}_L, \text{SIVD}_R$ are mean-error pairs of real numbers encoding the theoretical uncertainty ranges on the left and right boundaries
 - $\text{SIVD}^{\text{syst}}$ is computed based on geometrical considerations, in the manner illustrated in Fig. 2.7(b)

2.C.4 Algorithm usage

This section explain how the formalism presented throughout this document is effectively used for producing the results shown in Sections 2.3 and 2.4.

First, the formalism is used for producing the plots showing the $\text{SIVD}(\beta)$ dependence (“Model fitting” section). For either PG or MPG, for a specific k value and a specific β on-grid value, the drawn model SIVD uncertainty range (Fig. 2.3) is obtained after the following computational steps:

- 1: $(\alpha_L, \alpha_R, \alpha_{\text{fit}}, \alpha_{\text{err}}) = \text{FIT-1}(\beta, k)$ ▷ executing 1st-level fitting
- 2: $\tilde{\beta} = (\beta, \alpha_L, \alpha_R, \alpha_{\text{fit}}, \alpha_{\text{err}})$ ▷ creating composite fitting information

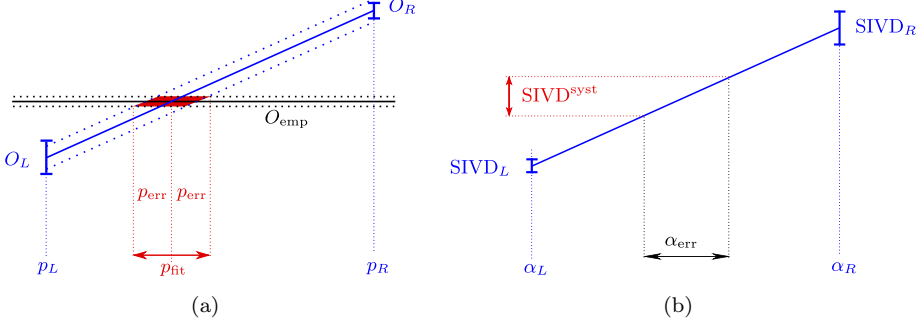


Figure 2.7: Illustration of computation carried out by INTERFITLIN (a) and by COMPSYSTEMERR (b), with the output quantities highlighted in red.

- 3: $\text{SIVD} = \text{NUMSIVD}(\bar{\beta}, k)$ ▷ numeric SIVD calculations
- 4: $(\text{SIVD}_{\text{avg}}, \text{SIVD}_{\text{stat}}, \text{SIVD}_{\text{syst}}) = \text{MEANSTATSYST}(\bar{\beta}, \text{SIVD})$

which provides the values of the SIVD average SIVD_{avg} , the SIVD statistical error $\text{SIVD}_{\text{stat}}$ and the SIVD systematic error $\text{SIVD}_{\text{syst}}$. One can then place a point at coordinates $(\beta, \text{SIVD}_{\text{avg}})$, within the respective k curve, with an error bar given by the maximum between $\text{SIVD}_{\text{stat}}$ and $\text{SIVD}_{\text{syst}}$.

Second, the formalism is used for providing the best-fitting, on-grid values for the α and β model parameters, which are used for generating sets of cultural vectors on which the LTCD-STCB analysis is applied (Sec. 2.4). For either PG or MPG and for a specific k value, the following procedure is followed:

- 1: $(\bar{\beta}_L, \bar{\beta}_R, \beta_{\text{fit}}, \beta_{\text{err}}) = \text{FIT-2}(k)$ ▷ Executing 2nd-level fitting
- 2: $(\beta_L, \alpha_L^L, \alpha_L^R, \alpha_L^{\text{fit}}, \alpha_L^{\text{err}}) = \bar{\beta}_L$ ▷ Decompressing left- β composite fitting information
- 3: $(\beta_R, \alpha_R^L, \alpha_R^R, \alpha_R^{\text{fit}}, \alpha_R^{\text{err}}) = \bar{\beta}_R$ ▷ Decompressing right- β composite fitting information
- 4: **if** $\beta_{\text{fit}} - \beta_L < \beta_R - \beta_{\text{fit}}$ **then** ▷ choosing β_L , since it is closer to β
- 5: $\beta = \beta_L$
- 6: **if** $\alpha_L^{\text{fit}} - \alpha_L^L < \alpha_L^R - \alpha_L^{\text{fit}}$ **then** ▷ choosing α_L^L , since it is closer to α_L^{fit}
- 7: $\alpha = \alpha_L^L$
- 8: **else** ▷ choosing α_L^R , since it is closer to α_L^{fit}
- 9: $\alpha = \alpha_L^R$
- 10: **end if**
- 11: **else** ▷ choosing β_R , since it is closer to β
- 12: $\beta = \beta_R$
- 13: **if** $\alpha_R^{\text{fit}} - \alpha_R^L < \alpha_R^R - \alpha_R^{\text{fit}}$ **then** ▷ choosing α_R^L , since it is closer to α_R^{fit}
- 14: $\alpha = \alpha_R^L$
- 15: **else**

```

16:          $\alpha = \alpha_R^R$                                  $\triangleright$  choosing  $\alpha_R^R$ , since it is closer to  $\alpha_R^{\text{fit}}$ 
17:     end if
18: end if

```

which provides the best on-grid values for the (α, β) pair.

Bibliography

- [1] Claudio Castellano, Santo Fortunato, and Vittorio Loreto. Statistical physics of social dynamics. *Rev. Mod. Phys.*, 81:591–646, May 2009.
- [2] Robert Axelrod. The dissemination of culture. *Journal of Conflict Resolution*, 41(2):203–226, 1997.
- [3] Konstantin Klemm, Víctor M. Eguíluz, Raúl Toral, and Maxi San Miguel. Global culture: A noise-induced transition in finite systems. *Phys. Rev. E*, 67:045101, Apr 2003.
- [4] Konstantin Klemm, Víctor M. Eguíluz, Raúl Toral, and Maxi San Miguel. Nonequilibrium transitions in complex networks: A model of social interaction. *Phys. Rev. E*, 67:026120, Feb 2003.
- [5] Marcelo N. Kuperman. Cultural propagation on social networks. *Phys. Rev. E*, 73:046139, Apr 2006.
- [6] Andreas Flache and Michael W. Macy. Local convergence and global diversity: The robustness of cultural homophily. *arXiv:physics/0701333*, 2007.
- [7] J. C. González-Avella, M. G. Cosenza, and K. Tucci. Nonequilibrium transition induced by mass media in a model for social influence. *Phys. Rev. E*, 72:065102, Dec 2005.
- [8] Damon Centola, Juan Carlos González-Avella, Víctor M. Eguíluz, and Maxi San Miguel. Homophily, cultural drift, and the co-evolution of cultural groups. *Journal of Conflict Resolution*, 51(6):905–929, 2007.
- [9] Jens Pfau, Michael Kirley, and Yoshihisa Kashima. The co-evolution of cultures, social network communities, and agent locations in an extension of Axelrod’s model of cultural dissemination. *Physica A: Statistical Mechanics and its Applications*, 392(2):381–391, 2013.
- [10] Federico Battiston, Vincenzo Nicosia, Vito Latora, and Maxi San Miguel. Robust multiculturalism emerges from layered social influence. *Sci. Rep.*, 7(1809), 2017.
- [11] Alex Stivala, Yoshihisa Kashima, and Michael Kirley. Culture and cooperation in a spatial public goods game. *Phys. Rev. E*, 94:032303, Sep 2016.

- [12] Luca Valori, Francesco Picciolo, Agnes Allansdottir, and Diego Garlaschelli. Reconciling long-term cultural diversity and short-term collective social behavior. *Proc. Natl. Acad. Sci.*, 109(4):1068–1073, 2012.
- [13] Alex Stivala, Garry Robins, Yoshihisa Kashima, and Michael Kirley. Ultrametric distribution of culture vectors in an extended Axelrod model of cultural dissemination. *Sci. Rep.*, 4(4870), 2014.
- [14] Alexandru-Ionuț Băbeanu, Leandros Talman, and Diego Garlaschelli. Signs of universality in the structure of culture. *The European Physical Journal B*, 90(12):237, Dec 2017.
- [15] Rama Cont and Jean-Philippe Bouchaud. Herd behavior and aggregate fluctuations in financial markets. *Macroeconomic Dynamics*, 4(02):170–196, 2000.
- [16] Michael Thompson, Richard J. Ellis, and Aaron Wildavsky. *Cultural Theory*. Westview Press, 1990.
- [17] Alan Page Fiske. The four elementary forms of sociality: framework for a unified theory of social relations. *Psychol Rev*, 99(4):689–723, 1992.
- [18] Harry C. Triandis. The psychological measurement of cultural syndromes. *Am Psychol*, 51(4):407–415, 1996.
- [19] Richard A. Shweder, Nancy C. Much, Manamohan Mahapatra, and Lawrence Park. The “big three” of morality (autonomy, community, divinity) and the “big three” explanations of suffering. In Allan M. Brandt and Paul Rozin, editors, *Morality and Health*, pages 119–169. Routledge, New York, 1997.
- [20] Jesse Graham, Brian A. Nosek, Jonathan Haidt, Ravi Iyer, Spassena Koleva, and Peter H. Ditto. Mapping the moral domain. *J Pers Soc Psychol*, 101(2):366–385, 2011.
- [21] Marco Verweij. Towards a theory of constrained relativism: Comparing and combining the work of pierre bourdieu, mary douglas and michael thompson, and alan fiske. *Sociological Research Online*, 12:1–15, 2007.
- [22] Joshua R. Bruce. Uniting theories of morality, religion, and social interaction: Grid-group cultural theory, the “big three” ethics, and moral foundations theory. *Psychology and Society*, 5:37–50, 2013.
- [23] Marco Verweij, Marieke van Egmond, Ulrich Kühnen, Shenghua Luan, Steven Ney, and M. Aenne Schoop. I disagree, therefore i am: how to test and strengthen cultural versatility. *Innovation: The European Journal of Social Science Research*, 27:1–12, 2014.
- [24] Steve Rayner. Cultural theory and risk analysis. In S. Krimsky and D. Golding, editors, *Social Theories of Risk*, pages 83–115. Praeger, 1992.

- [25] James Tansey and Steve Rayner. Cultural theory and risk. In Robert L. Heath and H. Dan O’Hair, editors, *Handbook of Risk and Crisis Communication*, pages 53–79. Routledge, New York, 2009.
- [26] Julia Poncela-Casasnovas, Mario Gutiérrez-Roig, Carlos Gracia-Lázaro, Julian Vicens, Jesús Gómez-Gardeñes, Josep Perelló, Yamir Moreno, Jordi Duch, and Angel Sánchez. Humans display a reduced set of consistent behavioral phenotypes in dyadic games. *Science Advances*, 2(8), 2016.
- [27] Jon Elster. *The Multiple Self*. Cambridge University Press, 1985.
- [28] Ernst Fehr and Karla Hoff. Introduction: Tastes, castes and culture: the influence of society on preferences. *The Economic Journal*, 121(556):F396–F412, 2011.
- [29] Alain Cohn, Ernst Fehr, and M. A. Marechal. Business culture and dishonesty in the banking industry. *Nature*, 516(7529):86–89, 2014.
- [30] A. Cohn, J Engelmann, E. Fehr, and Michael André. Maréchal. Evidence for countercyclical risk aversion: An experiment with financial professionals. *American Economic Journal*, 105(2):860–85, 2015.
- [31] Shalom H. Schwartz. Universals in the content and structure of values: Theoretical advances and empirical tests in 20 countries. volume 25 of *Advances in Experimental Social Psychology*, pages 1 – 65. Academic Press, 1992.
- [32] Richard Nisbett. *The Multiple Self*. The Free Press, 2003.
- [33] Ulrich Kühnen and Daphna Oyserman. Thinking about the self influences thinking in general: cognitive consequences of salient self-concept. *Journal of Experimental Social Psychology*, 38(5):492 – 499, 2002.
- [34] Karlheinz Reif and Anna Melich. Euro-barometer 38.1: Consumer protection and perceptions of science and technology, november 1992. <http://www.icpsr.umich.edu/icpsrweb/ICPSR/studies/06045>, 1995.
- [35] Jonathan L. Gross and Steve Rayner. *Measuring Culture: A Paradigm for the Analysis of Social Organization*. Columbia Univ Press, 1985.
- [36] Maroussia Favre and Didier Sornette. A generic model of dyadic social relationships. *PLoS ONE*, 10(3):1–16, 2015.
- [37] Vudtiwat Ngampruetikorn and Greg J. Stephens. Bias, belief, and consensus: Collective opinion formation on fluctuating networks. *Phys. Rev. E*, 94:052312, Nov 2016.
- [38] Peter Taylor. Algorithm for generating integer partitions. <http://math.stackexchange.com/questions/18659/algorithm-for-generating-integer-partitions>, 01 2011.

Chapter 3

Ultrametricity increases the predictability of cultural dynamics

A quantitative understanding of societies requires useful combinations of empirical data and mathematical models. Models of cultural dynamics aim at explaining the emergence of culturally homogeneous groups through social influence. Traditionally, the initial cultural traits of individuals are chosen uniformly at random, the emphasis being on characterizing the model outcomes that are independent of these ('annealed') initial conditions. Here, motivated by an increasing interest in forecasting social behavior in the real world, we reverse the point of view and focus on the effect of specific ('quenched') initial conditions, including those obtained from real data, on the final cultural state. We study the predictability, rigorously defined in an information-theoretic sense, of the *social content* of the final cultural groups (i.e. who ends up in which group) from the knowledge of the initial cultural traits. We find that, as compared to random and shuffled initial conditions, the hierarchical ultrametric-like organization of empirical cultural states significantly increases the predictability of the final social content by largely confining cultural convergence within the lower levels of the hierarchy. Moreover, predictability correlates with the compatibility of short-term social coordination and long-term cultural diversity, a property that has been recently found to be strong and robust in empirical data. We also introduce a null model generating initial conditions that retain the ultrametric representation of real data. Using this ultrametric model, predictability is highly enhanced with respect to the random and shuffled cases, confirming the usefulness of the empirical hierarchical organization of culture for forecasting the outcome of social influence models.

This chapter is based on the following scientific article:
A. I. Băbeanu, J. van de Vis and D. Garlaschelli, arXiv:1712.05959 (2017).

3.1 Introduction

Understanding the self-organization and emergence of large-scale patterns in real societies is one of the most fascinating, yet extremely challenging problems of modern social science [1]. A prominent field of research studies the spontaneous emergence of groups of culturally homogeneous individuals. One of the mechanisms that are believed to play a key role in this process is *social influence*, i.e. the gradual convergence of the cultural traits, attitudes and opinions of individuals subject to mutual social interactions. Stylized models of cultural dynamics under social influence have attracted the interest of an interdisciplinary community of sociologists, computational social scientists and statistical physicists [2].

One of the prototypical models in this context is the popular Axelrod model [3], which has been studied in many variants over the last two decades [4, 5, 6, 7, 8, 9, 10, 11, 12]. The model is multi-agent, with a cultural vector associated to each agent. One cultural vector is a sequence of subjective cultural traits (opinions, preferences, beliefs) that each agent possesses, with respect to a predefined set of features (variables, topics, issues). The dynamics is driven by social influence, which iteratively increases the similarity of the cultural vectors of pairs of interacting individuals. However, interactions are only allowed among pairs of individuals whose vectors are already closer than a certain (implicit or explicit) threshold distance, a mechanism known as *bounded confidence* and having its origins in the so-called ‘assimilation-contrast theory’ [13] in social science. The intuition behind the model, successfully confirmed via numerical simulations and analytic calculations, is that social influence increases cultural similarity, yet full convergence is precluded by bounded confidence. The net result is the emergence of a certain number of *cultural domains*, each containing several individuals with identical cultural vectors and mutually separated by a distance larger than the bounded confidence threshold, thus no longer interacting with each other. The value of the model is the identification of a viable, decentralized mechanism according to which cultural diversity can persist at a global (inter-domain) scale, even if it vanishes at a local (intra-domain) scale.

Given the focus on the qualitative aspect of such an emergent pattern, the Axelrod model has been traditionally studied by specifying uniformly random initial conditions for the cultural vectors of all individuals, i.e. by drawing each cultural trait independently from a probability distribution that is flat over the set of possible realizations. Consistently with this uninformative (and deliberately unrealistic) choice, the focus of many studies has been the characterization of the outcomes of the model that are robust upon averaging over multiple realizations of the initial randomness. Since the cultural dynamics evolving the initial state is also stochastic, a second average over the dynamics is also required. We may therefore say that this is the ‘annealed’ version of the model. Examples of quantities that are stable across multiple realizations of uniformly random initial conditions are the expected *number* and expected *size* of final cultural domains. An obvious counter-example is the *values* of the vectors ending up in such do-

mains: as follows from the complete symmetry in cultural space implied by the uniformity of the initial randomness, such values are by construction maximally unpredictable.

On the other hand, recent studies have investigated the model starting from different classes of initial conditions, beyond the uniformly random one. In particular, emphasis has been put on using initial conditions constructed from empirical data [14, 15, 16] (Chap. 1) and their randomized, trait-shuffled counterparts – obtained by randomly shuffling, for each component of the cultural vectors, the empirical values (traits) of all individuals in the sample. These studies have emphasized a strong dependence of the final outcome on the initial conditions. For instance, certain model outcomes that have an interesting interpretation in terms of enabling the coexistence of short-term social collective behavior and long-term cultural diversity [14] (more details are provided later in this paper) are found to vary significantly across the classes of empirical, trait-shuffled, and uniformly random initial conditions, while remaining largely stable when considering different instances belonging to the same class. This stability implies that empirical cultural data share certain remarkably universal properties, independent of the specific sample considered and at the same time significantly different from those exhibited by random and randomized data [16] (Chap. 1). This has stimulated the introduction of stochastic, structural models aimed at capturing the essential properties of the empirical cultural data [15, 17] (Chap. 2).

Strong dependence of cultural dynamics on the initial conditions might be a useful property to exploit in the light of the increasing interest towards forecasting social and cultural behavior in the real world. Examples include the predictability of certain aspects of political elections, public campaigns, spreading of (fake) news, financial bubbles and crashes, and commercial success of new items. If interest is shifted towards the predictability of future long-term outcomes given certain initial conditions, then a corresponding change of perspective is implied at the level of modeling. In particular, the aforementioned ‘annealed’ framework, where the outcome of models of cultural dynamics is averaged over multiple realizations of the initial randomness, becomes less relevant. On the contrary, if a specific (e.g. empirical) initial condition is known, it becomes natural to use it as the single initial specification of the heterogeneity of the system. Obviously, averaging with respect to different random trajectories of the social influence dynamics, all starting from the same initial cultural state, remains important and necessary. We may therefore call this the ‘quenched’ version of the model.

In this work we focus for the first time on the predictability of the *social content* of the cultural domains in the final state of the Axelrod model, given a certain initial state. By social content we mean the composition of the different domains in terms of individuals, i.e. we are interested in forecasting ‘who ends up in which cultural domain’. It should be noted that the social content is one of those properties that, just like the values of the final cultural vectors, is maximally unpredictable when considering the usual annealed model under uniformly random initial conditions. By contrast, we consider the quenched scenario start-

ing from specific initial conditions sampled from empirical, shuffled, random, and an additional, ‘ultrametric’ class of initial conditions.

We find that, remarkably, empirical and random initial conditions are associated with the highest and, respectively, lowest degree of predictability, which we rigorously define in an information-theoretic sense. This means that, as compared with the usual uniform specification of the initial conditions of the model, empirical data allow for a much more reliable forecast of the identity of the individuals forming the final cultural domains. We find that this result follows from the fact that the hierarchical, ultrametric-like organization of empirical cultural vectors, when coupled with bounded confidence, largely confines cultural convergence within the lower levels of the hierarchy. This result is confirmed using surrogate data that, while retaining only the ultrametric representation of real data, are also found to be associated with a higher predictability with respect to the shuffled and random conditions. The predictability associated to random and randomized cultural vectors is lower because it is difficult to identify a meaningful and robust hierarchical structure within the lower levels of which social influence remains confined.

Even if we do not perform an explicit analysis of the *cultural content* of the final domains, the finding that their social content is predictable, coupled with the fact that the initial cultural vectors of all individuals are known, implies that each final cultural vector will be a mixture of the traits of the initial vectors of the individuals ending up in the same cultural domain. This means that, the higher the predictability of the social content, the higher that of the cultural content as well. The take-home message is that the empirical hierarchical organization of culture and its ultrametric representation are very informative and useful for forecasting the outcome of models of cultural dynamics.

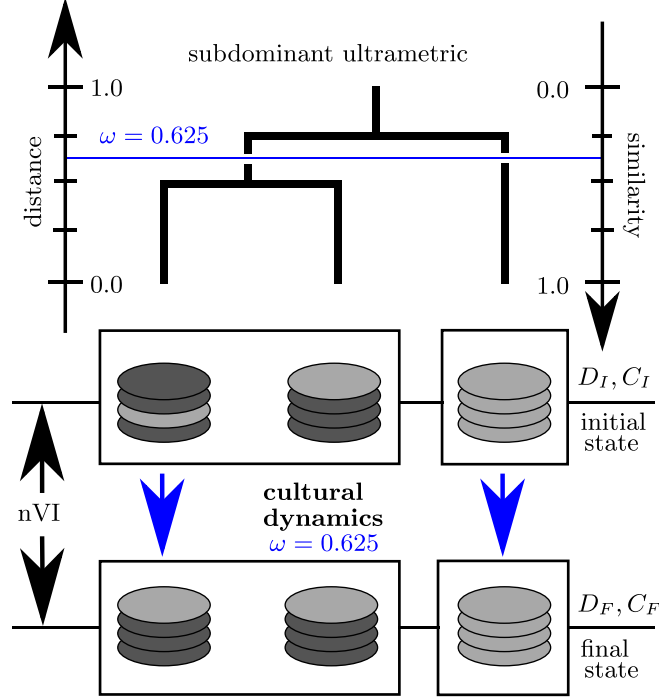


Figure 3.1: Cultural dynamics with an ultrametric initial state. At the top, a dendrogram with three leaves is shown, with a distance (or dissimilarity) scale on the left, with an associated similarity scale on the right and a threshold of $\omega = 0.625$ applied with respect to the former. The dendrogram is a subdominant ultrametric representation of distances between three cultural vectors, which are illustrated below its branches. These vectors are defined in terms of four binary variables (features), corresponding to the four horizontal rows of disks, whose possible values (traits) are denoted by the light-gray and dark-gray colors. The boxes show the initial state partition, formed by two clusters (and connected components) obtained by applying the $\omega = 0.625$ cut in the dendrogram. Together, the three vectors make up an initial cultural state on which the cultural dynamics model can be applied. For a bounded confidence value is set to $\omega = 0.625$, one of the possible final states is shown at the bottom. The boxes show the final state partition, formed by two cultural domains, within which cultural vectors are identical. The discrepancy between the initial state and final state partitions is measured with the normalized variation of information quantity nVI, which in this situation would give a value of 0.0, since the two partitions are identical.

3.2 Ultrametricity and cultural dynamics

The notion of ultrametricity refers to sets of objects that are hierarchically organized in certain abstract spaces, with applications in various fields, including mathematics (p -adic numbers), evolutionary biology (phylogenetic trees) and statistical physics (spin glasses) [18]. In practice, an ultrametric representation can be produced as the output of a hierarchical clustering algorithm applied to a matrix of pairwise distances between objects [18]. For the purpose of this work, these objects are the cultural vectors, whose pairwise cultural distances are computed in the same manner as in Refs. [16, 17, 15, 14] (also Chaps. 1 and 2) – the following explanations concerning ultrametricity are mostly restricted to cultural vectors, although many of the concepts have a wide range of applicability. The ultrametric representation of N cultural vectors can be visualized as a dendrogram (a binary hierarchical tree; see the top of Fig. 3.1) with N leaves (one for each vector) and $N - 1$ branching points (often referred to as “branchings”, for simplicity), sorted by $N - 1$ real numbers that are attached to them. These numbers can be defined in two, equivalent ways: on a distance scale (top-left axis) or on a similarity scale (top-right axis) – both quantities take values between 0.0 and 1.0, while adding up to 1.0. Each number is an approximation for distances between leaves that are first merged at the respective branching point. These $N - 1$ numbers and the topology of the dendrogram retain part of the information inherent in the cultural distance matrix (which is specified by $N(N - 1)/2$ numbers), so the dendrogram is an approximation of this matrix. The approximation is exact and algorithm-independent only when the original distances are perfectly ultrametric: a stronger version of the triangle inequality is satisfied for all triplets of distinct objects [18]. A cut can be performed at a certain height ω in the dendrogram, providing an ω -dependent partition of the N cultural vectors (see Fig. 3.1). For a dendrogram obtained via the single-linkage hierarchical clustering algorithm (See Ref. [19] and references therein), the ω -dependent partition is the same as that encoding the connected components obtained by applying an ω -threshold to the initial matrix of distances.

Ref. [14] pointed out that a dendrogram approximating an empirical cultural state shows a clearer hierarchical organization than those approximating its shuffled or random counterparts, suggesting that the ultrametric representation is better suited for empirical data than for shuffled or random data. In addition, cultural dynamics applied to the empirical cultural state appeared to mostly induce convergence within the clusters of the ω -dependent partition, if ω is equal to the bounded confidence threshold used in the cultural dynamics model (see below). These observations were made in a qualitative way, by visually inspecting dendrograms obtained with the average-linkage hierarchical clustering algorithm [20, 21]. Instead, we perform here a systematic, quantitative comparison between ω -dependent partitions of initial cultural states and associated partitions of final states resulting from cultural dynamics, for different classes of initial cultural states. In addition, one of these classes is defined by enforcing, on average,

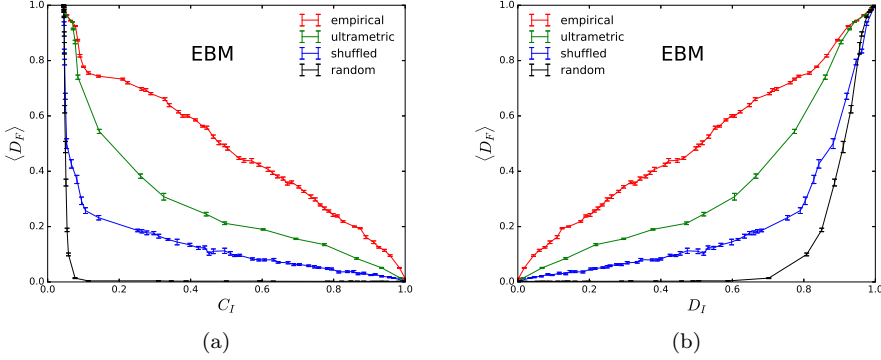


Figure 3.2: Relationships between the important diversity and coordination measures. One sees the dependence of the final, average diversity $\langle D_F \rangle$, first (a) on the initial coordination C_I , second (b) on the initial diversity measure D_I . This is shown for one empirical (red), one ultrametric-generated (green), one shuffled (blue) and one random (black) set of cultural vectors. All sets of cultural vectors have $N = 500$ elements and are defined with respect to the same cultural space, from the variables of the empirical Eurobarometer (EBM) data. The errors of $\langle D_F \rangle$ are standard mean errors obtained from 10 cultural dynamics runs.

the ultrametric representation of empirical data, generalizing a method originally proposed in Ref. [22] for biological taxonomies. Whenever an ultrametric representation is constructed within this study, the single-linkage algorithm [19] is used instead of the average-linkage one, since it provides the subdominant ultrametric, which is the ‘closest from below’ to the original distances and unique [23], while also being equivalent to the hierarchical connected-component representation, as mentioned above. This choice is also common for the purpose of evaluating measures of ultrametricity, like the cophenetic correlation coefficient, which is done in Ref. [15].

Cultural dynamics is modeled here by a simple, Axelrod-type model, without any underlying geometry for a social network or a geographical-physical space: essentially, all N agents are connected to each other. Instead, a bounded-confidence threshold ω is present, controlling the maximum cultural distance for which social influence can successfully occur. This is exactly the model used in Refs. [14, 16, 17] (also in Chaps. 1 and 2) and partly in Ref. [15]. As anticipated in Sec. 3.1, this model converges to a random final, absorbing state, one that consists of domains of internally identical and externally non-interacting cultural vectors – distances within such groups are zero, while distances across are larger or equal to ω .

Fig. 3.1 captures the essence of this study. At the center, the figure shows

an initial cultural state with 3 vectors, defined in terms of 4 binary features, with possible traits (values) denoted by the two shades of gray. Each of the three vectors is matched to a branch of the dendrogram drawn at the top, which encodes the subdominant ultrametric representation of the initial cultural state. For this specific case, the distance between the first two vectors is 0.5, while the distances between any of these two and the third are 0.75, which together make up a perfectly ultrametric discrete space, thus exactly matching the distances encoded by the dendrogram. The horizontal line denotes a possible ω -cut that can be applied to the dendrogram, which induces a splitting into two (in the example shown) branches and two associated subsets of vectors, which together form a ω -dependent partition (or clustering) of the initial set. This partition is the same as that induced by the set of connected cultural components of the ω -thresholded cultural graph. At the bottom, the figure shows one possible final state resulting from the cultural dynamics process, for a bounded confidence threshold set to the same ω value as the dendrogram cut. The groups of identical vectors constitute another, ω -dependent partition characterizing the cultural state, which exactly matches, in this case, the initial state partition. Other final configurations are possible, due to the stochastic nature of cultural dynamics. It is even possible, although unlikely, that by a succession of convenient interactions the second vector “migrates” from the cluster on the left to the one on the right during the dynamics. The abundance of such deviations is quantitatively studied below, for several classes of initial conditions.

3.3 Partition-specific quantities

The initial and final partitions form the basis of all calculations performed in this study. Each type of partition is characterized by two types of quantities, denoted by (D_I, C_I) for initial partitions and by (D_F, C_F) for final partitions. These quantities are referred to as the coordination (C_I and C_F) and the diversity measures (D_I and D_F). They are computed according to the following formulas:

$$D_a(\omega) = \frac{N_C^a(\omega)}{N}, \quad C_a(\omega) = \sqrt{\sum_A \left(\frac{S_A^a}{N} \right)^2}, \quad (3.1)$$

where $a \in \{I, F\}$ distinguishes between “initial” and “final”, N_C^a is the number of clusters (connected components if $a = I$, groups of identical vectors if $a = F$), and S_A^a is the size of cluster A for the given ω value. Note that D_a is a measure of diversification, while C_a is a measure of non-homogeneity encoded by the respective partition. Moreover, since cultural dynamics is a stochastic process, it is meaningful to talk about averages over final state partitions (over multiple dynamical runs), which is particularly useful for the final diversity measure $\langle D_F(\omega) \rangle$.

The $\langle D_F(\omega) \rangle$ quantity has been interpreted as a measure of propensity to long-term cultural diversity, while the $C_I(\omega)$ has been interpreted as a measure of propensity to short-term collective behavior [14, 16] (Chap. 1). Through their common dependence on ω , the correspondence between the two quantities is graphically illustrated in Fig. 3.2(a). Along each curve, different points correspond to different ω values, while different curves correspond to different classes of initial conditions. It is clear that the empirical cultural state allows for much more compatibility between the aspects measured by the two quantities than the shuffled and the random cultural state, as pointed out in Ref. [14]. In fact, this is the analysis used in Ref. [14] to highlight the structure of empirical cultural data and in Ref. [16] (Chap. 1) to emphasize the universality of this structure – except for the “ultrametric” scenario, which is first introduced here. In this scenario, a set of N cultural vectors is generated such that, on average, the pairwise distances reproduce those encoded in the subdominant ultrametric representation of an empirical set of cultural vectors of the same N . This is achieved using an extension of the method developed in Ref. [22], in the context of genetic sequences. The extension here allows the method to work with combinations of features of different ranges and types, where the range stands for the number of traits and the type indicates whether the feature is ordinal or nominal. This is described in detail in Appendix 3.A. On the other hand, a shuffled set of cultural vectors is obtained by randomly and independently permuting empirical cultural traits among vectors, with respect to every feature, thus exactly enforcing the empirical trait frequencies. Note that the ultrametric cultural state comes closer to the empirical behavior than the shuffled cultural state, suggesting that empirical ultrametric is better than empirical trait frequencies at explaining the generic empirical structure. Finally, a random set of cultural vectors is obtained by drawing each trait at random, from a uniform probability distribution, while only retaining the empirical data format, and thus the cultural space – determined the number of features, together with the range and type of each feature. Eurobarometer 38.1 [24] data is used here, formatted according to the procedure in Ref. [16] (Chap. 1).

For the same four sets of cultural vectors used in Fig. 3.2(a), the average final diversity $\langle D_F(\omega) \rangle$ is plotted against the initial diversity $D_I(\omega)$ in Fig. 3.2(b). This visualization, previously used [14, 15] without the ultrametric scenario, illustrates the extent to which cultural dynamics preserves the number of clusters when going from the initial to the final partition. As observed before, the number of clusters is well preserved by cultural dynamics acting on empirical data, which happens much less for shuffled data and even less for random data. This goes along with the idea that the final partition can be predicted from the initial partition if empirical data is used for specifying the latter. Note that, like in Fig. 3.2(a), ultrametric-generated data lies in between the empirical and shuffled scenarios, confirming that the subdominant ultrametric information, which is directly related to the sequence of ω -dependent initial partitions, is rather robust with respect to cultural dynamics.

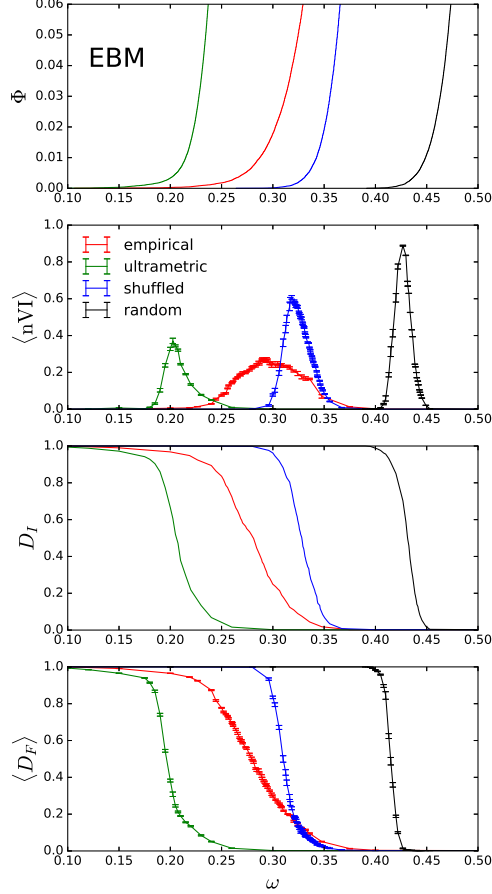


Figure 3.3: Visualization of the ultrametric predictability of cultural dynamics. The dependence on the bounded-confidence threshold ω is shown for several quantities: most importantly, the normalized variation of information between the initial and final partitions $\langle nVI \rangle$ at the center-top; the fraction of initially active cultural links Φ at the top; the initial diversity D_I at the center-bottom; the final, average diversity $\langle D_F \rangle$ at the bottom. This is shown for one empirical (red), one ultrametric-generated (green), one shuffled (blue) and one random (black) set of cultural vectors. All sets of cultural vectors have $N = 500$ elements and are defined with respect to the same cultural space, from the variables of the Eurobarometer (EBM) data. The errors of $\langle D_F \rangle$ and $\langle nVI \rangle$ are standard mean errors obtained from 10 cultural dynamics runs.

3.4 Predictability of the final state

Although informative, the comparison between the $\langle D_F(\omega) \rangle$ and $D_I(\omega)$ is incomplete as a way of assessing the predictability of the final partition from the initial partition: two partitions might have the same number of clusters, but the sizes and/or contents of these clusters might be very different. In order to take all this into account in a consistent way, the discrepancy between the initial and final state partitions is evaluated using the variation of information measure VI [25], as a function of ω . This is an information-theoretic measure that acts as a metric distance within the space of possible partitions of a set of N elements. It is convenient to work with the normalized version of this quantity $\text{nVI}(\omega) = \text{VI}(\omega) / \log(N)$, which retains the meaning and metricity of the original quantity, as long as N remains the same ($N = 500$ for all results presented here).

The dependence of $\langle \text{nVI} \rangle$ on ω is shown in the second panel of Fig. 3.3, for the same 4 cultural states used in Fig. 3.2, where the averaging is performed over multiple dynamical runs, like for the $\langle D_F \rangle$ quantity. The empirical state shows the lowest maximal $\langle \text{nVI} \rangle$ value, followed by the ultrametric, the shuffled and the random states. This figure shows, in a rigorous way, that the outcome of cultural dynamics can be predicted relatively well based on the initial state, if this is constructed from empirical data and comparably well if this is constructed based on the empirical ultrametric information. On the other hand, shuffled and random data exhibit lower predictability. Note that, for either scenario, $\langle \text{nVI} \rangle$ vanishes for the low- ω and the high- ω regions, which is where both the initial and final partitions consist of N single-object clusters and of one, N -objects cluster respectively. This can be understood by looking at the dependence of the D_I and $\langle D_F \rangle$ quantities on ω shown in the in the third and fourth panels: the ω region for which $\langle \text{nVI} \rangle$ is significantly larger than 0.0 is roughly the region where either D_I or $\langle D_F \rangle$ is substantially different from 1.0 and 0.0.

In parallel, the first panel of Fig. 3.3 shows the ω -dependence of the fraction of initially active cultural links Φ : the fraction of pairs (i, j) of cultural vectors whose distance $d_{ij} < \omega$ in the initial state. This shows that the ω interval that is non-trivial with respect to D_I , $\langle D_F \rangle$ and $\langle \text{nVI} \rangle$ seems to be largely determined by the shape of Φ , which is nothing else than the cumulative distribution of inter-vector distances. The properties of this distribution – average lower for empirical data than for random data, standard deviation higher for empirical data than for either shuffled or random data – have been studied before [14, 15] and are recognizable in the first panel of Fig. 3.3. Note that, for the ultrametric scenario, the interesting ω region and the Φ profile are compressed in a lower- ω region compared to empirical data. This means that the branchings in the dendrogram obtained from ultrametric-generated data occur at lower ω values than those in the dendrogram obtained from the original, empirical data. In turn, this is due to the distances between the ultrametric-generated cultural vectors reproducing, on average, the subdominant ultrametric empirical distances, rather than the original empirical distances, while the former are known to systematically underestimate

the latter, particularly for higher distance values, as long as the empirical vectors are not perfectly ultrametric, which in practice is always the case.

There is another aspect that can be noted when comparing, for either scenario, the shape of $\Phi(\omega)$ in the first panel with the shape of $D_I(\omega)$ in the third panel of Fig. 3.3: as ω is decreased, most of the cultural links need to be eliminated in order to reach the abrupt region of the $D_I(\omega)$ transition, for which the number of clusters in the initial partition becomes comparable to N . This is not surprising on general grounds. For instance, the Erdős-Rényi model of random graphs [26] exhibits a critical link density of $1/N$, at which a giant connected component is present, if N is the number of nodes in the graph, instead of the number of cultural vectors. Still, this analogy should not be taken too far. The random graph interpretation is closest to the random cultural state scenario used here, since the expected pairwise distance entailed by the latter is the same for any pair of cultural vectors, just like the connection probability entailed by the former is the same for any pair of nodes. However, even the random scenario has an underlying metric structure, due to how cultural spaces are defined[16] (Chap. 1), which should introduce more triangles than expected otherwise, while the shuffled and empirical scenarios are additionally affected by inhomogeneities in their cultural space distributions.

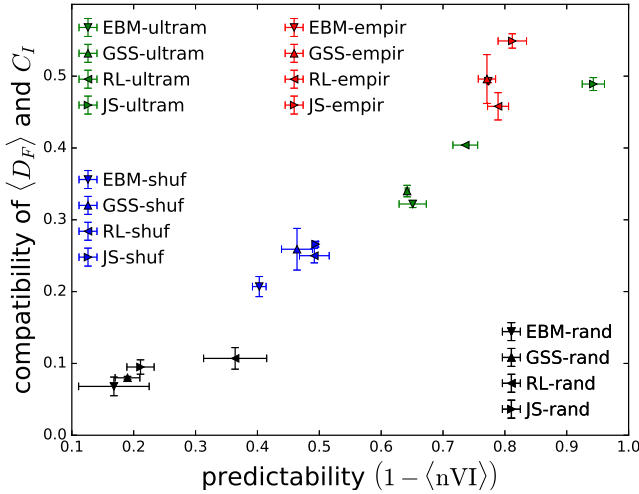


Figure 3.4: Relationship between compatibility of final diversity and initial coordination (vertical axis) and predictability of the final partition from the initial partition. Each point corresponds to one cultural state, belonging to one class and to one empirical source: each color corresponds to one class of cultural states, while marker type correspond to one dataset, as indicated in the legends. All cultural state consist of $N = 500$ cultural vectors.

The analysis presented in Figs. 3.2 and Fig. 3.3 was repeated for three other datasets: the General Social Survey [27], Jester[28] and the Religious Landscape [29], processed according to the formatting rules of Ref. [16] (Chap. 1). For all four datasets, the results are presented in a joint, compact manner by means of Fig. 3.4, while more detailed results are shown in Appendix 3.B. Each of the points in the figure corresponds to a combination of one dataset and one scenario. The vertical axis corresponds to a measure of compatibility between long-term cultural diversity $\langle D_F \rangle$ and short-term collective behavior C_I , namely a measure of the overall departure of the $\langle D_F \rangle$ vs C_I curve from the lower-left corner in Fig. 3.2(a). The horizontal axis corresponds to a measure of predictability of the final state from the initial state, namely an inverse measure of the overall departure of the $\langle nVI \rangle$ vs ω from the horizontal axis in the second panel of Fig. 3.3.

For both measures, simple definitions are employed: rather than integrating information from every ω value for which some departure is present, both definitions conceptually rely only on one, representative ω^* point, for which both departures are relatively high. Specifically, ω^* is defined by intersecting the $\langle D_F \rangle$ vs C_I curve with the main diagonal $\langle D_F \rangle = C_I$. In practice, since just a finite number of ω values are available for any combination of dataset and scenario, one uses instead the two ω values that are closest to the main diagonal of the $\langle D_F \rangle$ vs C_I plot from either of the two sides. These two values, labeled as ω_L and ω_R , “bracket” ω^* from the left and right respectively: $\omega_L < \omega^* < \omega_R$. The ω^* itself is never explicitly calculated, but is conceptually useful for the explanations below.

The compatibility approximates the distance between the $(\langle D_F(\omega^*) \rangle$ vs $C_I(\omega^*))$ point and the $(\langle D_F \rangle = 0, C_I = 0)$ point, normalized by the length of the main diagonal of the $\langle D_F \rangle$ vs C_I plot. In practice, this is evaluated in terms of ω_L and ω_R according to:

$$\frac{\sqrt{\langle D_F(\omega_L) \rangle^2 + C_I^2(\omega_L)} + \sqrt{\langle D_F(\omega_R) \rangle^2 + C_I^2(\omega_R)}}{2\sqrt{2}},$$

while the associated error is evaluated as:

$$\frac{\sqrt{\langle D_F(\omega_L) \rangle^2 + C_I^2(\omega_L)} - \sqrt{\langle D_F(\omega_R) \rangle^2 + C_I^2(\omega_R)}}{2\sqrt{2}}.$$

The predictability approximates the distance between the $(\omega^*, \langle nVI(\omega^*) \rangle)$ point and the $\langle nVI \rangle = 1$ line. In practice, this is evaluated as:

$$1 - \frac{\langle nVI(\omega_L) \rangle + \langle nVI(\omega_R) \rangle}{2},$$

while the associated error is evaluated as:

$$\frac{|\langle nVI(\omega_L) \rangle - \langle nVI(\omega_R) \rangle|}{2}.$$

Note that compatibility increases with predictability in a roughly linear way, at least for the cultural states considered here. Moreover, cultural states belonging to the same class tend to cluster together in the compatibility-predictability space. A notable exception is ultrametric-Jester, which is significantly outside the ultrametric class in terms of predictability, showing higher predictability than any of the empirical states. Still, it is clear that cultural states that are closer to the universal $\langle D_F \rangle$ vs C_I empirical behavior also allow for better estimates of the final partition from the initial one.

The observed increase of compatibility with predictability provides some insights about the nature of empirical data, or at least about the shape of an empirical-like dendrogram characteristic for the upper-right corner of Fig. 3.4. This can be understood by realizing that the ultrametric and empirical states approach an ideal, limiting situation of perfect predictability, for which the initial and final partitions are identical irrespective of ω . This implies that $\langle D_F(\omega) \rangle = D_I(\omega)$ and consequently that the $\langle D_F \rangle$ vs C_I curve is essentially the D_I vs C_I curve and thus controlled by the geometry of the subdominant ultrametric dendrogram. One can then show – see Appendix Sec. 3.C – that this geometry needs to be highly “unbalanced” in order to explain the close-to-linear $\langle D_F \rangle \approx 1 - C_I$ empirical behavior in Fig. 3.2(a) and the compatibility values of approximately 0.5 following from it. For a perfectly-unbalanced geometry, the k th highest dendrogram branching separates only one leaf from the remaining $N - k$, for all $k \in \{1, \dots, N - 1\}$. By contrast, a perfectly-balanced geometry entails a splitting into two, equal clusters for each dendrogram branching, which would induce an inverse square $\langle D_F \rangle \propto C_I^{-2}$ behavior – see Appendix Sec. 3.C – closer to that of shuffled and random cultural states, with a lower compatibility value. Thus, while going from the random to the empirical class, by enforcing more and better empirical information, the increasing level of compatibility becomes more suggestive of an unbalanced dendrogram geometry, while the increasing level of predictability increases the reliability of this geometric interpretation.

3.5 Conclusion

This study focused on the ultrametric representation of sets of cultural vectors used for specifying the initial state of cultural dynamics models. On one hand, it introduced another procedure for randomly generating initial conditions based on the subdominant ultrametric information of empirical data. On the other hand, it examined the extent to which the subdominant ultrametric representation can be used for predicting the final state of cultural dynamics in a simple theoretical setting. The bounded-confidence threshold parameterising the dynamical model was used to extract an initial-state partition from the ultrametric representation. This was systematically compared, in terms of variation of information, with the corresponding final state partition consisting of groups of identical cultural vectors. The comparison showed that the predictive power of the ultramet-

ric is relatively high for empirical cultural states, which are closely followed by ultrametric-generated states, which are followed by the shuffled and then by the random states. Moreover, higher predictability appears to go hand in hand with higher compatibility between a propensity to long-term cultural diversity and a propensity to short-term collective behaviour, which was previously shown to be a hallmark of empirical structure. This means that ultrametric information is better than trait-frequency information at explaining this structure. These results further advance the understanding of the relationship between ultrametricity and cultural dynamics. Moreover, it is tempting to speculate that, for the purpose of forecasting the dynamics of culture in the real world, knowledge about the current distribution of individuals in cultural space might be sufficient, with little or no need for running simulations, at least if one assumes that consensus-favoring social influence is the essential driving force of this dynamics.

Appendices

3.A Ultrametric-generation method

This section explains the method for generating sets of cultural vectors belonging to the “ultrametric” class. The method is an extension of that developed in Ref. [22]. The description here is somewhat similar to that in Ref. [22], but the nomenclature specific to cultural vectors is used, instead of that specific to genetic sequences.

The method takes as input a dendrogram, as well as a target cultural space – the number of cultural features F , together with the range (number of traits) q and type (nominal or ordinal) of each feature. This information is taken from empirical data and the single-linkage hierarchical clustering algorithm is employed for constructing the dendrogram whenever the method is used in this study. Upon every use, the method generates, in a stochastic way, a set of N cultural vectors associated to the N leaves of the dendrogram, such that, on average, the pair-wise similarities between cultural vectors match the similarities encoded by the dendrogram.

More precisely, for each cultural feature in the target space, the method enforces:

$$E[s_{ij}^q] = \rho_{\alpha_{ij}}, \quad (3.2)$$

where $E[...]$ stands for “expectation value”, α_{ij} is the lowest branching in the dendrogram joining leaves i and j , $\rho_{\alpha_{ij}}$ is the similarity encoded by this branching and s_{ij}^q is the partial contribution to the similarity between cultural vectors i and

j of a feature of range q , which is computed according to the following formula:

$$s_{ij}^q = \begin{cases} \delta(x_i^k, x_j^k) & \text{if nominal,} \\ 1 - \frac{|x_i^k - x_j^k|}{q^k - 1} & \text{if ordinal,} \end{cases} \quad (3.3)$$

which depends on whether the feature is nominal or ordinal, where δ stands for the Kronecker delta function, x_i^k and x_j^k are the traits of vectors i and j with respect to feature k with range q^k – for ordinal features k , the traits are marked with integers between 1 to q^k . Eq. (3.3) is consistent with the cultural distance definition in Refs. [14, 15, 16, 17] (and Chaps. 1 and 2) – as mentined above: similarity = 1.0 – distance.

In Eq. (3.2), the expectation $E[\dots]$ implies averaging over multiple runs of the method, for the same dendrogram and the same cultural feature. Although in practice the method is used only once (and independently) for each feature, the fact that a large number F of features are present makes this approach sensible: the expectation $E[s_{ij}]$ of the complete similarity s_{ij} will also match $\rho_{\alpha_{ij}}$ (since the complete similarity is the arithmetic average of the feature-level similarities), while the fluctuations of s_{ij} around $\rho_{\alpha_{ij}}$ will decrease with F . In other words, as pointed out in Ref. [22], the expectation in Eq. (3.2) can be interpreted in two idealized ways: averaging over infinitely many runs or averaging over infinitely many features.

In order to enforce Eq. (3.2) for every pair (i, j) , the method controls for the extent to which the traits of different vectors are choosen independently of each other. For every feature, all the N choosen cultural traits originate in independent random draws from a uniform probability distribution, but the number of draws is smaller or equal to N . Thus, the traits of vectors i and j either originate in the same draw, with probability P_{ij} , or originate in different draws, with probability $1 - P_{ij}$. In the former case the two traits are identical, with a well-determined feature-level similarity $s_{ij}^q = 1$. In the latter case, the two traits may be identical or different, so that s_{ij}^q fluctuates around an expectation value $f(q)$. Taking both cases into account, the expectation value of s_{ij}^q is:

$$E[s_{ij}^q] = P_{ij} + [1 - P_{ij}]f(q), \quad (3.4)$$

where the expectation for different draws $f(q)$ reads:

$$f(q) = \begin{cases} \frac{1}{q} & \text{if nominal,} \\ \frac{2q-1}{3q} & \text{if ordinal,} \end{cases} \quad (3.5)$$

which is the expression of the expected, feature-level similarity between two traits drawn at random from a uniform probability distribution, obtained analytically from Eq. (3.3) for either type of features. The choices of traits and the associated random draws are mangaged by the stochastic-algorithmic part of the method (briefly explained at the end of this section), which is designed to ensure that:

$$P_{ij} = \rho_{\alpha_{ij}}^I \quad (3.6)$$

is satisfied, where $\rho_{\alpha_{ij}}^I$ is a corrected version of the similarity $\rho_{\alpha_{ij}}$ implicit in the α_{ij} branching:

$$\rho_{\alpha_{ij}}^I = \rho_{\alpha_{ij}} - h(\rho_{\alpha_{ij}}, q), \quad (3.7)$$

where h is a correction function chosen such that Eqs. (3.2) holds, subject to (3.4) and (3.6). Specifically, by combining Eq. (3.6) with Eq. (3.4) and then with Eq. (3.2), one obtains:

$$\rho_{\alpha_{ij}}^I + [1 - \rho_{\alpha_{ij}}^I]f(q) = \rho_{\alpha_{ij}}. \quad (3.8)$$

By inserting Eq. (3.7) in Eq. (3.8) and further manipulations, one obtains the following expression for the correction function:

$$h(\rho_{\alpha}, q) = \frac{1 - \rho_{\alpha}}{1 - f(q)} f(q). \quad (3.9)$$

Note that Eq. (3.6) identifies $\rho_{\alpha_{ij}}^I$ with a probability, meaning that $\rho_{\alpha}^I > 0$ should be satisfied for all branchings α . This implies, given Eq. (3.7) and Eq. (3.9), that $\rho_{\alpha} > f(q)$ for all branchings α of the given dendrogram and for all features in the target space. This condition needs to be satisfied in order for this method to be valid and is actually satisfied by all four empirical dendrograms used in this study. Also note that the method in Ref. [22] is recovered as a special case of the above, by restricting to nominal features of constant q via Eq. (3.5).

Finally, it is worth describing the stochastic-algorithmic part of the method. For each of the F features in the target space, the following steps are carried out:

- the dendrogram is recursively explored starting with the root branching; for every branching α reached by this exploration, one of the following two things happens:
 - one of the q traits is randomly chosen, according to a uniform distribution and assigned to all cultural vectors corresponding to leaves under branching α , without further exploring any branching below α ;
 - the exploration is continued with each of the two branches emerging from α , if that branch leads to another branching, instead of leading to a leaf;

with probability Q_{α} for the former and probability $1 - Q_{\alpha}$ for the latter, where:

$$Q_{\alpha} = \frac{\rho_{\alpha}^I - \rho_{p(\alpha)}^I}{1 - \rho_{p(\alpha)}^I}, \quad (3.10)$$

where $p(\alpha)$ is the parent branching of α , if α is not the root, while $\rho_{p(\alpha)}^I = 0$ if α is the root.

- for each of the leaves whose traits are not assigned during the above step, one of the q traits is randomly chosen, according to a uniform distribution and assigned to the respective cultural vector.

This algorithmic procedure ensures that Eq. 3.6 holds, for reasons that are fully explained in Ref. [22].

It is worth noting that the ultrametric-generation method described in this section makes use of all the information inherent in the geometry of the dendrogram that it receives as input – both the topology and the similarities ρ encoded by the branching points of the dendrograms are used. However, the generated sets of cultural vectors will in general not be precisely ultrametric, in the strict mathematical sense [18] (unless it is applied in the limit of F being much larger than N). Still, they are generated based on the empirical ultrametric information and are arguably as close as they can be to reproducing the ultrametric set of pairwise distances.

3.B Detailed results

This section shows the complete results concerning the ω -dependence of relevant quantities, for the other three data sets that are used in this study in addition to the Eurobarometer (EBM [24]): the General Social Survey (GSS [27]) data in Fig. 3.5, the Religious Landscape (RL [29]) data in Fig. 3.6 and the Jester (JS [28]) data in Fig. 3.7. Each of these three figures follows the format of Fig. 3.3 above, with four panels and four scenarios. Although, for each type of scenario, there is a certain variability in the width and location of the non-trivial ω interval, the results are qualitatively similar to those obtained for EBM data, with a notable exception visible for the analysis of Jester data in Fig. 3.7: the second panel shows that the discrepancy between the initial and the final partition, as measured by $\langle nVI \rangle$, is clearly smaller for the ultrametric cultural state than for the empirical cultural state, so the overall predictability is higher. This is in agreement with the observation made in relation to Fig. 3.4 about the relatively high predictability value of the Jester-ultrametric point.

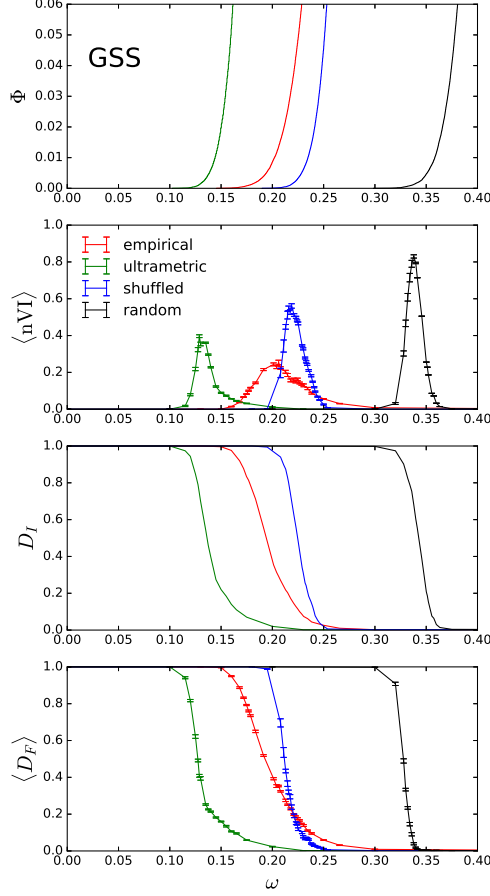


Figure 3.5: Visualization of the ultrametric predictability of cultural dynamics. The dependence on the bounded-confidence threshold ω is shown for several quantities: most importantly, the normalized variation of information between the initial and final partitions $\langle nVI \rangle$ at the center-top; the fraction of initially active cultural links Φ at the top; the initial diversity D_I at the center-bottom; the final, average diversity $\langle D_F \rangle$ at the bottom. This is shown for one empirical (red), one ultrametric-generated (green), one shuffled (blue) and one random (black) set of cultural vectors. All sets of cultural vectors have $N = 500$ elements and are defined with respect to the same cultural space, from the variables of the General Social Survey (GSS) data. The errors of $\langle D_F \rangle$ and $\langle nVI \rangle$ are standard mean errors obtained from 10 cultural dynamics runs.

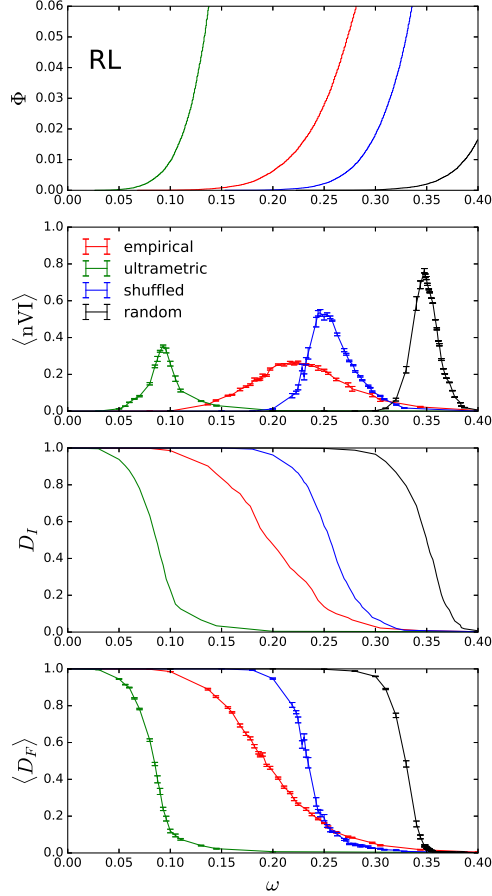


Figure 3.6: Vizualisation of the ultrametric predictability of cultural dynamics. The dependence on the bounded-confidence threshold ω is shown for several quantities: most importantly, the normalized variation of information between the initial and final partitions $\langle nVI \rangle$ at the center-top; the fraction of initially active cultural links Φ at the top; the initial diversity D_I at the center-bottom; the final, average diversity $\langle D_F \rangle$ at the bottom. This is shown for one empirical (red), one ultrametric-generated (green), one shuffled (blue) and one random (black) set of cultural vectors. All sets of cultural vectors have $N = 500$ elements and are defined with respect to the same cultural space, from the variables of the Religious Landscape (RL) data. The errors of $\langle D_F \rangle$ and $\langle nVI \rangle$ are standard mean errors obtained from 10 cultural dynamics runs.

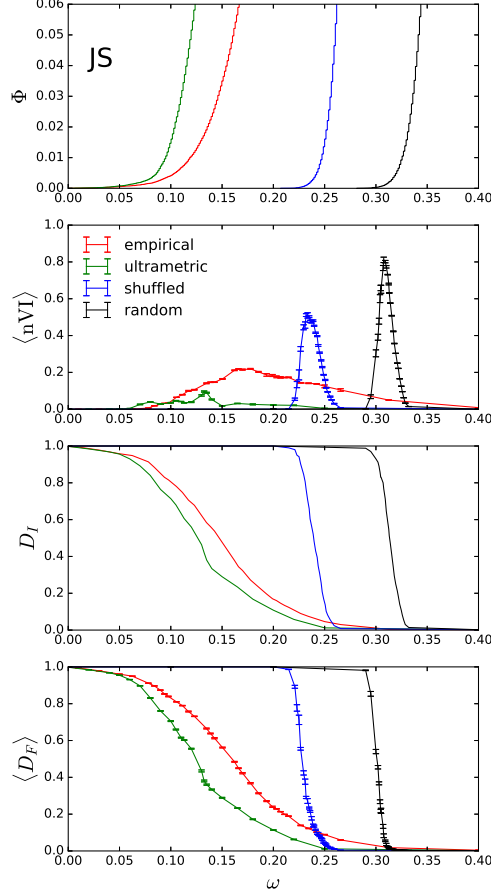


Figure 3.7: Visualization of the ultrametric predictability of cultural dynamics. The dependence on the bounded-confidence threshold ω is shown for several quantities: most importantly, the normalized variation of information between the initial and final partitions $\langle nVI \rangle$ at the center-top; the fraction of initially active cultural links Φ at the top; the initial diversity D_I at the center-bottom; the final, average diversity $\langle D_F \rangle$ at the bottom. This is shown for one empirical (red), one ultrametric-generated (green), one shuffled (blue) and one random (black) set of cultural vectors. All sets of cultural vectors have $N = 500$ elements and are defined with respect to the same cultural space, from the variables of the Jester (JS) data. The errors of $\langle D_F \rangle$ and $\langle nVI \rangle$ are standard mean errors obtained from 10 cultural dynamics runs.

3.C Dendrogram geometry

This section gives some analytical insight on how the dendrogram geometry is related to the behaviour of the two measures of initial diversity D_I and initial coordination C_I . As functions of ω , the two measures only change (in steps) when ω crosses the distance value associated to any of the branchings of the dendrogram. Thus, one can replace the dependence of D_I and C_I on ω with a dependence on k , which counts the number of dendrogram branchings above a given ω , in terms of their associated distance values – k increases from 0 to $N - 1$ as ω decreases from 1.0 to 0.0. Based on Eq. (3.1), one can thus write:

$$D_I(k) = \frac{N_C^I(k)}{N}, \quad C_I(k) = \sqrt{\sum_A \left(\frac{S_A^I}{N} \right)_k^2}. \quad (3.11)$$

There are two, extreme types of dendrogram geometries that are worth considering, the "perfectly-unbalanced geometry" and the "perfectly-balanced geometry". These are illustrated in Fig. 3.8.

For the perfectly-unbalanced geometry, shown on the left side of Fig. 3.8, the number of connected components is:

$$N_C^I(k) = k + 1, \quad (3.12)$$

while the sizes of the connected component are:

$$S_A^I(k) = \begin{cases} N - k, & \text{if } A = 1 \\ 1, & \text{if } A \in \{2, 3, \dots, k + 1\} \end{cases}. \quad (3.13)$$

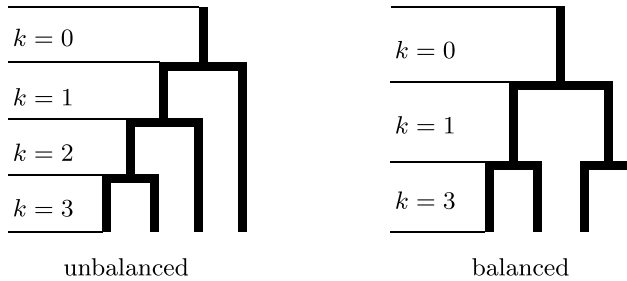


Figure 3.8: Sketch of a “perfectly balanced” (left) dendrogram geometry and a “perfectly unbalanced” (right) one, for $N = 4$ leaves. The values of k indicate the number of branchings above any cut that would be applied to the dendrogram within the respective horizontal band.

From Eqs. (3.11) and (3.12), one obtains the behaviour of the initial diversity measure:

$$D_I(k) = \frac{k+1}{N}, \quad (3.14)$$

while from Eqs. (3.11) and (3.13) one obtains the behaviour of the initial coordination measure:

$$C_I(k) = \sqrt{\left(\frac{N-k}{N}\right)^2 + k\left(\frac{1}{N}\right)^2}, \quad (3.15)$$

from which it follows that:

$$C_I(k) = \sqrt{1 - 2\frac{k}{N} + \frac{k^2}{N^2} + \frac{k}{N^2}}, \quad (3.16)$$

where one can neglect the $\frac{k}{N^2}$ term in the limit of large N , thus obtaining:

$$C_I(k) \approx 1 - \frac{k}{N}. \quad (3.17)$$

From Eqs. 3.14 and 3.17 it follows that:

$$C_I(k) \approx 1 - D_I(k) - \frac{1}{N}, \quad (3.18)$$

which can be rephrased, after neglecting the $\frac{1}{N}$ term in the limit of large N , to:

$$D_I(k) \approx 1 - C_I(k), \quad (3.19)$$

which describes the second-diagonal empirical behaviour of Fig. 3.2(a), under the assumption that $D_F(k) = D_I(k), \forall k$.

For a perfectly-balanced geometry, shown on the right side of Fig. 3.8, the only relevant values of k (those corresponding to non-vanishing ω intervals) are $k = \sum_{i=0}^{l-1} 2^i$, with $l \in \{0, 1, 2, \dots, \log_2 N\}$. For these values of k , the number of connected components, like in the unbalanced case, is described by Eqs. (3.12), while the sizes of the connected components are:

$$S_A^I(k) = N/(k+1), \forall A \in \{1, 2, \dots, k+1\}, \quad (3.20)$$

from which it follows that the initial coordination measure is:

$$C_I(k) = \sqrt{(k+1)\left(\frac{1}{k+1}\right)^2} = \frac{1}{\sqrt{k+1}}. \quad (3.21)$$

Since the k -dependence of the initial diversity measure D_I , like in the unbalanced case, is described by Eq. (3.14), it follows that:

$$D_I(k) = \frac{1}{NC_I^2(k)}, \quad (3.22)$$

which, under the assumption that $D_F(k) = D_I(k), \forall k$, entails a curve more similar to that of the shuffled or random curves of Fig. 3.2(a), than to that of the empirical curve. Moreover, this curve comes arbitrarily close to the lower left corner as N increases.

To sum up, the above reasoning shows that, as long as $D_F(\omega) = D_I(\omega), \forall \omega$, an unbalanced dendrogram geometry fits the empirical $D_F(C_I)$ behaviour very well, while a balanced dendrogram geometry does not. Although the latter entails a $D_F \propto C_I^{-2}$ behaviour quite similar to that observed for shuffled or random data, one cannot say that a balanced geometry is a good description for either of these two cases, since the assumption that $D_F = D_I$ is false for both these cases, for the interesting ω intervals.

Bibliography

- [1] Mark Buchanan. *The Social Atom*. Bloomsbury, New York, NY, 2007.
- [2] Claudio Castellano, Santo Fortunato, and Vittorio Loreto. Statistical physics of social dynamics. *Rev. Mod. Phys.*, 81:591–646, May 2009.
- [3] Robert Axelrod. The dissemination of culture. *Journal of Conflict Resolution*, 41(2):203–226, 1997.
- [4] Konstantin Klemm, Víctor M. Eguíluz, Raúl Toral, and Maxi San Miguel. Global culture: A noise-induced transition in finite systems. *Phys. Rev. E*, 67:045101, Apr 2003.
- [5] Konstantin Klemm, Víctor M. Eguíluz, Raúl Toral, and Maxi San Miguel. Nonequilibrium transitions in complex networks: A model of social interaction. *Phys. Rev. E*, 67:026120, Feb 2003.
- [6] Marcelo N. Kuperman. Cultural propagation on social networks. *Phys. Rev. E*, 73:046139, Apr 2006.
- [7] Andreas Flache and Michael W. Macy. Local convergence and global diversity: The robustness of cultural homophily. *arXiv:physics/0701333*, 2007.
- [8] J. C. González-Avella, M. G. Cosenza, and K. Tucci. Nonequilibrium transition induced by mass media in a model for social influence. *Phys. Rev. E*, 72:065102, Dec 2005.
- [9] Damon Centola, Juan Carlos González-Avella, Víctor M. Eguíluz, and Maxi San Miguel. Homophily, cultural drift, and the co-evolution of cultural groups. *Journal of Conflict Resolution*, 51(6):905–929, 2007.
- [10] Jens Pfau, Michael Kirley, and Yoshihisa Kashima. The co-evolution of cultures, social network communities, and agent locations in an extension of

- Axelrod's model of cultural dissemination . *Physica A: Statistical Mechanics and its Applications*, 392(2):381–391, 2013.
- [11] Federico Battiston, Vincenzo Nicosia, Vito Latora, and Maxi San Miguel. Robust multiculturalism emerges from layered social influence. *Sci. Rep.*, 7(1809), 2017.
 - [12] Alex Stivala, Yoshihisa Kashima, and Michael Kirley. Culture and cooperation in a spatial public goods game. *Phys. Rev. E*, 94:032303, Sep 2016.
 - [13] Muzafer Sherif and Carl Iver Hovland. *Social judgment: Assimilation and Contrast Effects in Communication and Attitude Change*. Yale University Press, New Haven, CT, 1961.
 - [14] Luca Valori, Francesco Picciolo, Agnes Allansdottir, and Diego Garlaschelli. Reconciling long-term cultural diversity and short-term collective social behavior. *Proc. Natl. Acad. Sci.*, 109(4):1068–1073, 2012.
 - [15] Alex Stivala, Garry Robins, Yoshihisa Kashima, and Michael Kirley. Ultrametric distribution of culture vectors in an extended Axelrod model of cultural dissemination. *Sci. Rep.*, 4(4870), 2014.
 - [16] Alexandru-Ionuț Băbeanu, Leandros Talman, and Diego Garlaschelli. Signs of universality in the structure of culture. *The European Physical Journal B*, 90(12):237, Dec 2017.
 - [17] Alexandru-Ionuț Băbeanu and Diego Garlaschelli. Evidence for mixed rationalities in preference formation. *arXiv:1705.06904v2*, 2017.
 - [18] R. Rammal, G. Toulouse, and M. A. Virasoro. Ultrametricity for physicists. *Rev. Mod. Phys.*, 58:765–788, Jul 1986.
 - [19] R. Sibson. Slink: An optimally efficient algorithm for the single-link cluster method. *The Computer Journal*, 16(1):30, 1973.
 - [20] Michael R. Anderberg. Chapter 6 - hierarchical clustering methods. In Michael R. Anderberg, editor, *Cluster Analysis for Applications*, Probability and Mathematical Statistics: A Series of Monographs and Textbooks, pages 131 – 155. Academic Press, 1973.
 - [21] Robert Reuven Sokal and Charles Duncan Michener. A statistical method for evaluating systematic relationships. *University of Kansas Scientific Bulletin*, 28:1409–1438, 1958.
 - [22] M. Tumminello, F. Lillo, and R. N. Mantegna. Generation of hierarchically correlated multivariate symbolic sequences. *The European Physical Journal B*, 65(3):333–340, 2008.

- [23] Rammal, R., Angles d'Auriac, J.C., and Doucot, B. On the degree of ultrametricity. *J. Physique Lett.*, 46(20):945–952, 1985.
- [24] Karlheinz Reif and Anna Melich. Euro-barometer 38.1: Consumer protection and perceptions of science and technology, november 1992. <http://www.icpsr.umich.edu/icpsrweb/ICPSR/studies/06045>, 1995.
- [25] Marina Meilă. Comparing clusterings – an information based distance. *Journal of Multivariate Analysis*, 98(5):873 – 895, 2007.
- [26] Paul Erdős and Alfréd. Rényi. On random graphs, i. *Publicationes Mathematicae (Debrecen)*, 6:290–297, 1959.
- [27] Tom W. Smith, Peter Marsden, Michael Hout, and Jibum Kim. General social surveys, 1993 ed. <http://gss.norc.org/get-the-data/spss>, 1972-2012.
- [28] Ken Goldberg, Theresa Roeder, Dhruv Gupta, and Chris Perkins. Eigentaste: A constant time collaborative filtering algorithm. *Information Retrieval*, 4(2):133–151, 2001.
- [29] Luis Lugo, Sandra Stencel, John Green, and Gregory Smith et al. U.s. religious landscape survey. religious beliefs and practices: Diverse and politically relevant. <http://www.pewforum.org/2008/06/01/>, 2008.

Chapter 4

A random matrix perspective of cultural structure

Recent studies have highlighted interesting structural properties of empirical cultural states: collections of vectors of cultural traits of real individuals, based on which one defines matrices of similarities between individuals. This study provides further insights about the structure encoded in these states, using concepts from random matrix theory. For generating random matrices that are appropriate as a structureless reference, we propose a null model that enforces, on average, the empirical occurrence frequency of each possible trait. With respect to this null model, the empirical similarity matrices show deviating eigenvalues, which may be signatures of cultural groups that might not be recognizable by other means. However, they can conceivably also be artifacts of arbitrary, dataset-dependent correlations between cultural variables. In order to understand this possibility, independently of any empirical information, we study a toy model which explicitly enforces a specified level of correlation in a minimally-biased way, in the simplest conceivable setting. In parallel, a second toy model is used to explicitly enforce group structure, in a very similar setting. By analyzing and comparing cultural states generated with these toy models, we show that a deviating eigenvalue, such as those observed for empirical data, can also be induced by correlations alone. Such a “false” group mode can still be distinguished from a “true” one, by evaluating the uniformity of the entries of the respective eigenvector, while checking whether this uniformity is statistically compatible with the null model. For empirical data, the eigenvector uniformities of all deviating eigenvalues are shown to be compatible with the null model, suggesting that the apparent group structure is not genuine, although a decisive statement requires further research.

This chapter is based on the following scientific article:
A. I. Băbeanu, arXiv:1803.04324 (2018).

4.1 Introduction

Understanding the complex behavior of social systems greatly benefits from constructively combining the increasing amount of empirical data with a variety of quantitative, theoretical approaches, often originating in the natural sciences [1, 2]. Although much of this interdisciplinary research focuses on the network and connectivity aspects of social systems [3], efforts are also being made for understanding a complementary aspect: the formation and dynamics of opinions, preferences, attitudes and beliefs, more generally referred to as “cultural traits” [4]. In particular, recent studies have placed a stronger emphasis on using empirical data about the cultural traits of real individuals [5, 6, 7, 8, 9] (Chaps. 1, 2 and 3). Such data is typically recorded within a short period of time from a random sample of people in a population, via a social survey with a large number of questions, so that a vector (or sequence) of cultural traits can be constructed for every individual, where each trait is an answer to one of the questions. The collection of all cultural vectors constructed from one empirical source is called an empirical “cultural state”, or an empirical “set of cultural vectors”, since it can be used to empirically specify the initial conditions of an Axelrod-type model of cultural dynamics [10]. Using previously developed tools [5, 6] that relied on models of cultural and opinion dynamics, Ref. [7] (Chap. 1) showed that empirical cultural states are characterized by properties that are highly robust across different data sets. These properties have been further explored [8, 9] (Chaps. 2 and 3) but not entirely understood. This generic structure present in an empirical cultural state is largely retained by the person-person matrix of cultural similarities that can be defined based on the cultural vectors, allowing for this structure to be further investigated by means of a random matrix approach.

Random matrix theory [11, 12] has been successfully for a variety of applications, such as the analysis of financial systems [13]. The framework deals with the properties of random matrices, under certain distributional assumptions. The associated statistical ensembles of matrices are used to compute the expected values (or even the probability distributions) of interesting, matrix dependent quantities. These theoretical expectations can be compared to empirical counterparts evaluated on matrices that encode information about the real world systems that are being studied. Statistically significant deviations of the empirical quantities are then interpreted as interesting, non trivial structural properties of the respective systems. The focus is on the eigenvalue spectrum of the empirical matrix, which can be, for instance, a correlation matrix between the time series recording the price dynamics of stocks [14] or the activity neurons [15]. In such cases, the appropriate assumptions of randomness are captured by the the Marchenko-Pastur [16] law, which gives a limiting distribution for the spectrum. The eigenmodes whose eigenvalues are significantly larger than what is expected based on the Marchenko-Pastur law are interpreted as joint dynamical patterns in terms of which the non-trivial behavior of the system can be understood, while the other are interpreted as the noise components of the system. Recently, Ref. [17]

extended this approach to similarity matrices constructed from categorical data, where an entry of the matrix is a similarity between two time series of discrete symbols. For instance, for one of the data sets of Ref. [17], each sequence of symbols corresponds to an electoral constituency of India, with different symbols associated to different winning parties and successive time steps associated to successive elections.

In this study, this random matrix approach is applied to spectra of empirical matrices of cultural similarities, constructed from data previously used in Refs. [5, 6, 7, 8, 9] (Chaps. 1, 2 and 3). Instead of relying on analytic formulas for estimating and filtering the noise, we make extensive use of numerical methods. This allows for a detailed investigation of three null models (Sec. 4.2), among which the uniformly random generation is the simplest and conceptually closest [17] to the approach of Marchenko and Pastur. As a second null model, we make use of trait shuffling, which is known to be important for understanding empirical cultural states, independently from spectral decomposition and random matrix notions [5, 6, 7, 8, 9] (Chaps. 1, 2 and 3), since it reproduces exactly the empirical trait occurrence frequencies. We propose an additional null model which also reproduces these empirical trait frequencies on average, while also incorporating some mathematically desirable properties of the uniform random generation. We name this procedure "restricted random generation". These null models are compared in terms of how well they reproduce the upper boundary of the noisy spectral region ("the bulk"), as well as the position of the highest eigenvalue. This is a strong outlier which can be understood as a "global mode", which for similarity matrices is guaranteed to be present even under the uniformly random scenario [17]. As shown in Sec. 4.2, the restricted random generation turns out to be more appropriate and is thus selected for further analysis. Based on restricted randomness, we numerically evaluate the probability distribution of the upper noise boundary, showing that there are several empirical eigenvalues significantly above this boundary. These correspond to modes that capture the structure in empirical data, since they are incompatible with null hypothesis behind restricted randomness. Hence, this manuscript will often refer to them as "structural modes".

It is tempting to interpret these deviating eigenmodes as signatures of group structure, in a manner similar to time series analysis [14], in the sense that the individuals are organized in terms of several cultural groups or categories. This is particularly intriguing, given that Ref. [8] (Chap. 2) provides indirect evidence for cultural structure being governed by a small number of cultural prototypes supposedly induced by universal "rationalities". However, it is important to keep in mind that the empirical data also shows pairwise correlations between cultural variables, that are at least partly induced by arbitrary, dataset-dependent similarities between how the corresponding items are defined, as previously pointed out [5, 7, 6] (Chap. 1). Since these correlations are not retained by restricted randomness and shuffling, it is possible that deviating eigenmodes are a direct consequence of them. This rises the question of whether these eigenmodes are

signatures of authentic group structure or are just artifacts of arbitrary correlations between variables. First, we explicitly show that, at least in principle, one can differentiate between the “correlations scenario” and the “groups scenarios” – this is not trivial, since group structure also entails, to a certain extent, pairwise correlations between features. This is done in Sec. 4.3 by studying, in a highly simplified, abstract setting, consisting of only binary features, two probabilistic models for generating (sets of) cultural vectors. The first model, labeled “FCI” (Sec. 4.3.1), explicitly enforces a certain pairwise coupling between all features, in a manner that gives rise to a certain level of correlation, without introducing any additional assumption or bias in the underlying probability distribution. This is ensured by a maximum-entropy derivation [18], which leads to a statistical ensemble that is mathematically equivalent to the canonical ensemble of the Ising model on a fully-connected lattice [19], where each feature corresponds to one lattice site and each cultural vector corresponds to a spin configuration. The second model, labeled “S2G” (Sec. 4.3.2), explicitly enforces a dual group structure, whose strength can be analytically tuned to match the first model in terms of the level of feature-feature correlations that arise as a side effect. More details about these models and about the interpretation of deviating eigenmodes are given in Sec. 4.3.

For any given level of feature-feature correlations, the two models are used for generating sets of cultural vectors. The structure of the resulting similarity matrices is captured by the subleading eigenvalue and its eigenvector. However, Sec. 4.4 shows that the expected strength and significance of the subleading eigenvalue is exactly the same for the FCI and S2G models, for any given correlation level, so the subleading eigenvalue does not discriminate between the two scenarios, confirming that the presence of deviating eigenvalues does not necessarily imply the presence of group structure. We show that the essential difference between the FCI and S2G is captured by the entries of the eigenvector associated to the subleading eigenvalue. In particular, the uniformity of these entries, quantified by the “eigenvector entropy” (see Sec. 4.4), shows a clearly different behavior as a function of correlation for the two models, with S2G showing an increasingly higher uniformity as the correlation level increases. Moreover, the dependence of the second-highest eigenvector entropy on the correlation level reproduces well the symmetry-breaking phase transitions that characterize the two models. In each case, the eigenvector entropy suddenly jumps from a regime of compatibility to a regime of incompatibility with the null model, exactly when the probability distribution associated to the respective model becomes bi-modal. The critical correlation associated to this transition is almost ten times smaller for S2G than for FCI. This further justifies the use of eigenvector entropy as an indicator of group structure in empirical data, as a complement to eigenvalue information.

Along these lines, Sec. 4.5 presents an enhanced analysis of empirical data, showing how the eigenmodes are distributed in terms of eigenvalue and eigenvector entropy, in comparison to expectations based on restricted randomness. Interestingly, all the structural modes previously identified in empirical data (based

on eigenvalues) are actually compatible with the null model in terms of eigenvector uniformity (based on eigenvector entropy). This is the case for all three datasets used in this study, suggesting that structural modes of culture are actually artifacts of arbitrary correlations between cultural variables. However, such a conclusion is conditional on the S2G model introduced here being representative, in a qualitative way, for any type of authentic cultural groups that empirical data might capture. As explained in Sec. 4.6, this might actually not be the case, so the presence of groups cannot be entirely rejected based on this study, especially if these groups are highly entangled. In particular, the S2G model does not account for the “mixing”, “multiple-self” ingredient that has been shown to be crucial [8] (Chap. 2) for an interpretation of cultural structure in terms of a small number of prototypes [6], inspired by social science frameworks such as Plural Rationality Theory [20]. In fact, it appears likely that structural modes induced by mixing prototypes would not exhibit higher eigenvector uniformities than expected based on the null model, while they should still qualify as group modes. More research is needed to establish whether this is indeed the case and, if so, to find a way of distinguishing structural modes induced by mixing prototypes from those induced by correlations.

4.2 Eigenvalue distributions for empirical data and null models

In this section, the eigenvalue spectra of empirical matrices of cultural similarities are evaluated. At the same time, three null models are evaluated and compared. Each null model is used to numerically generate similarity matrices, by randomly sampling from the associated statistical ensemble, which enforces, to a certain extent, the empirical information that is expected to not be of interest – this is information that, on a priori grounds, clearly has more to do with arbitrary survey design choices than with any authentic cultural structure. One of these models, namely the “restricted random” model, which is first introduced here, is chosen as a good benchmark with respect to which interesting structure is to be measured, as explained below. Before presenting the actual results, some mathematical clarifications are given with respect to the computation of similarity matrices, the spectral decomposition procedure and the definitions of the null models.

A cultural similarity matrix is a square, $N \times N$ matrix obtained from N cultural vectors, which are all defined with respect the same set of F cultural features (variables or dimensions). Each feature can take one of q^k possible, discrete values, called “cultural traits”, where k labels the features, according to some order that is arbitrary, but consistent across all vectors. Moreover, each feature can be either nominal, marked as $f_{\text{nom}}^k = 1$, or ordinal, marked as $f_{\text{nom}}^k = 0$, which affects how its similarity contribution is defined. Each entry s_{ij} of the similarity

matrix is then computed according to:

$$s_{ij} = \frac{1}{F} \sum_{k=1}^F \left[f_{\text{nom}}^k \delta(x_i^k, x_j^k) + (1 - f_{\text{nom}}^k) \left(1 - \frac{|x_i^k - x_j^k|}{q^k - 1} \right) \right], \quad (4.1)$$

encoding the similarity between vectors i and j , where δ stands for the Kronecker delta function and x_i^k and x_j^k denote the traits recorded with respect feature k in vectors i and j respectively – for the ordinal case, it is important that x_i^k and x_j^k take discrete, rational values between 1 and q^k , while for the nominal case they only need to take symbolic values from any (feature-specific) alphabet. Note that the similarity measure in Eq. (4.1) is an arithmetic average of the similarity contributions of the F cultural features, in agreement with Refs. [5, 6, 7, 8, 9] (Chaps. 1, 2 and 3) – although in these studies most concepts are presented in terms of cultural distances d_{ij} , these have a trivial relationship to cultural similarities: $d_{ij} = 1 - s_{ij}$. For an empirical matrix, each vector i corresponds to one individual in the real world, each feature k to one question or item in the questionnaire used to collect the data, so that the realized trait x_i^k , which lies at the intersection between vector i and feature k , corresponds to the answer/rating given by individual i to question/item k . For a matrix generated based on a null model, the N vectors are generated according to the specified random procedure, while retaining (at least) the empirical data format, namely the type f_{nom}^k and range q^k of each feature k . Note that, in contrast to the empirical symbolic sequences used in Ref. [17], cultural vectors have no axis of time, so everything is equivalent up to a reordering of the cultural features, as long as this is done consistently for all cultural vectors. This is irrelevant for any of the mathematical operations involved by the analysis here, but it is relevant for the interpretation: cultural vectors capture no time-evolution, and should be interpreted as instantaneous, multidimensional opinion profiles, rather than as dynamical, one-dimensional dynamical profiles.

From Eq. (4.1) it follows that such a similarity matrix is real and symmetric, from which it follows, according to the spectral theorem, that it has N real eigenvalues with N associated orthonormal eigenvectors with real entries. This implies that the matrix can be decomposed in the following way:

$$s_{ij} = \sum_{l=1}^N \lambda_l v_l^i v_l^j, \quad (4.2)$$

where “ λ_l ” and “ v_l ” are used to denote the l th highest eigenvalue and, respectively, the eigenvector associated to it, while v_l^i is the i th entry of eigenvector v_l . Throughout this study, special attention is payed to λ_1 and λ_2 , the highest and second highest eigenvalues of various similarity matrices, also denoted as the “leading” and “subleading” eigenvalues respectively. In parallel, “ λ ” is used to denote any generic eigenvalue. More notation will be introduced below, as needed.

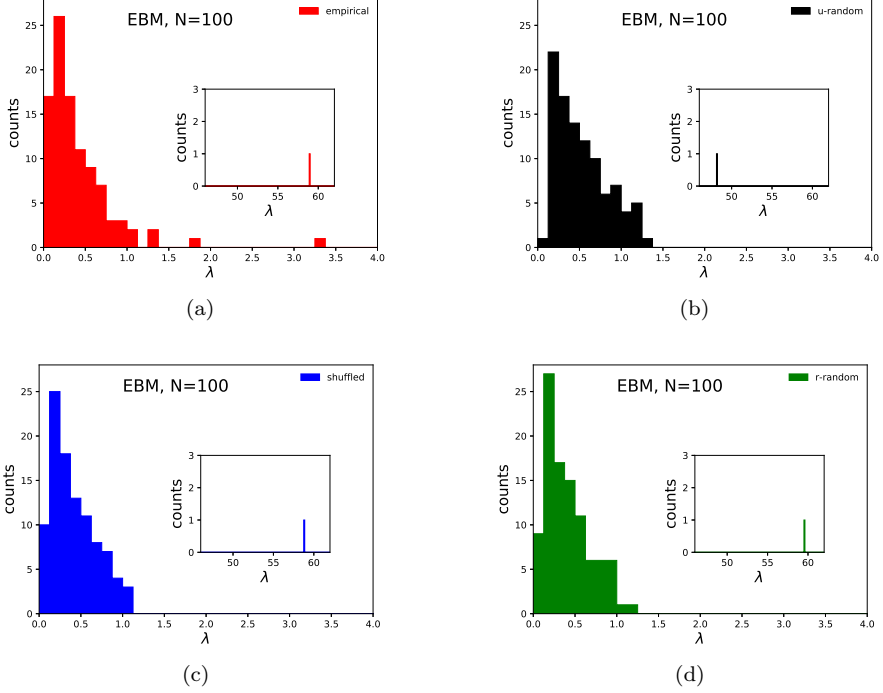


Figure 4.1: Eigenvalue spectra of cultural similarity matrices. The first panel correspond to an empirical cultural state (a) with $N = 100$ vectors constructed from Eurobarometer (EBM) data, while the other three correspond to associated cultural states generated with the uniform random (b), the shuffled (c) and the restricted random (d) null models, using partial information from the empirical cultural state and the same $N = 100$. For each panel, the inset shows the leading eigenvalue of the respective spectrum. For comparison purposes, the axis ranges and bin widths are the same across the four panels, for the main plots as well as for the insets.

All similarity matrices used in this study are based on sets of $N = 100$ cultural vectors, regardless of whether they are empirical, generated with one of the three null models introduced below or with one of the toy models introduced in Sec. 4.3. At the same time, the number of features F is always larger than N , which ensures that the information in every similarity matrix is not redundant, since its number of entries $N \times N$ is smaller than the number of entries in the set of cultural vectors $F \times N$.

Fig. 4.1(a) shows the eigenvalue spectrum of an empirical similarity matrix computed based on $N = 100$ cultural vectors extracted from Eurobarometer (EBM) data, which records attitudes and opinions of European Union citizens on various topics concerning technology, the environment and certain policy issues [21]. The data is formatted according to the procedure described in Ref. [7] (Chap. 1), which makes $F = 144$ cultural features available. The vertical axis gives the number of eigenvalues occurring in each bin along the horizontal axis. The inset focuses on the higher λ region of the horizontal axis, where the leading eigenvalue λ_1 is located. The high value of λ_1 is expected based on purely mathematical grounds [17], due to the overall positivity of any such similarity matrix. In most cases, all entries of the eigenvector associated to λ_1 have the same sign and very similar absolute values, meaning that, according to Eq. (4.2), the $\lambda_1 v_1^i v_1^j$ captures a large, highly uniform, positive component of the matrix entries s_{ij} . The λ_1 eigenmode thus accounts for the overall tendency towards similarity of the entire system, which is partly due to how similarity is defined and partly (see below) due to feature-level non-uniformities. For this reason, the λ_1 mode will also be referred as the “global mode”, term which originates from time-series analysis [14] based on correlation matrices, for which a global mode may or may not be present, depending on the system. Using exactly the same format as Fig. 4.1(a), each of the other three panels of Fig. 4.1 shows the spectrum of a similarity matrix generated from each of the three null models: “uniform randomness”, “shuffling” and “restricted randomness”.

First, Fig. 4.1(b) shows the spectrum of a similarity matrix generated via uniform randomness (abbreviated as “u-random”). Specifically, for every vector, each trait is chosen independently at random from the traits available at the level of the respective feature, with equal probability attached each possible trait. This means that uniform randomness retains minimal information from the empirical cultural state used for Fig. 4.1(a): only the number of features, the type and the number of traits of each feature. Note that the leading eigenvalue of this matrix is comparable to that of the empirical matrix. Ref. [17] showed that the analytic, limiting distribution given by the Marchenko-Pastur formula has a shape that is qualitatively similar to the bulk of the u-random spectrum. Quantitatively however, the analytic and numerical distributions become truly similar only if an important parameter controlling the former is left free and fit to the numerical results, instead of being directly set to F/N , which can be done when dealing with matrices of correlations between N time series with F numerical entries each. Moreover, the Marchenko-Pastur formula completely fails to describe the

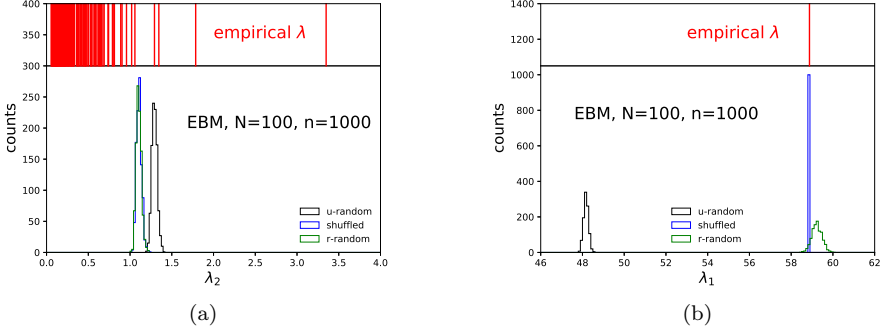


Figure 4.2: Leading and subleading eigenvalue distributions for random matrices. The figure shows the subleading eigenvalue λ_2 distribution (a) and the leading eigenvalue λ_1 distribution (b), for the three null models (legends), implementing uniform randomness (black), shuffling (blue) and restricted randomness (green), in comparison to the empirical eigenvalues, whose positions are marked by the vertical (red) lines in the upper bands. For each distribution, $n = 1000$ similarity matrices are numerically generated from the respective null model. Everything is based on the same set of $N = 100$ vectors constructed from Eurobarometer (EBM) data used in Fig. 4.1.

leading eigenvalue.

Second, Fig. 4.1(c) shows the eigenvalue spectrum of a similarity matrix generated via shuffling. Specifically, with respect to every feature, the traits realized in the empirical state are randomly permuted among the vectors, such that every permutation is equally likely. This is done independently for every feature, so that all types of correlations between features are destroyed. The procedure preserves exactly the number of times each trait is empirically realized, in addition to preserving the data format of the empirical state in Fig. 4.1(a). Note that, by construction, the assignment of traits to vectors is not entirely independent across vectors, implying that the number of vectors N resulting from shuffling has to be exactly the same as the number of empirical vectors used.

Third, Fig. 4.1(d) shows the spectrum of a similarity matrix generated via restricted randomness (abbreviated as “r-random”). Specifically, for every vector, each trait is chosen independently at random from the traits available at the level of the respective feature, with different probabilities attached to the possible traits, these probabilities being directly proportional to the empirical occurrence frequencies of the respective traits. This means that, like the shuffling procedure, restricted randomness also reproduces the empirical trait frequencies, but on average. Moreover, it also retains the independent generation specific to uniform randomness, which allows for an arbitrary number N of cultural vectors to be

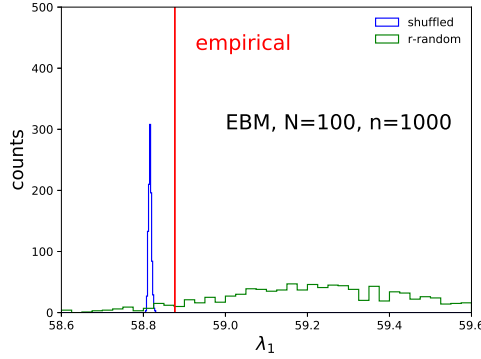


Figure 4.3: Detailed comparison in terms of the leading eigenvalue. The figure shows in detail the distributions of the leading eigenvalue λ_1 for the shuffled (blue) and restricted random (green) null models, in comparison to the empirical value (vertical red line). For each distribution, $n = 1000$ similarity matrices are numerically generated from the respective null model. Everything is based on the same set of $N = 100$ vectors constructed from Eurobarometer (EBM) data used in Fig. 4.1. For visual purposes, the bin size of the shuffled histogram is ten times smaller than for the restricted random histogram.

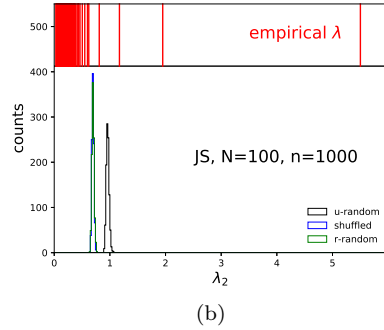
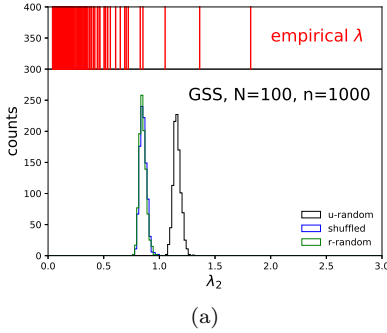


Figure 4.4: Empirical structure in other datasets. The figure shows the subleading eigenvalue λ_2 distribution for the three null models (legends), implementing uniform randomness (black), shuffling (blue) and restricted randomness (green), in comparison to the empirical eigenvalues, whose positions are marked by the vertical (red) lines in the upper band, based on General Social Survey data (a) and for Jester data (b). In each case, $N = 100$ cultural vectors are constructed from the respective dataset. For each null model, $n = 1000$ random matrices of $N = 100$ vectors are generated for drawing the associated distribution.

generated, regardless of how large this number is for empirical data. The independent generation should also make the analytic tractability of the model easier. Although neither of these two advantages are directly exploited in this study, they suggest that restricted randomness is conceptually more appropriate than either uniform randomness or shuffling, as it incorporates the desirable properties of both.

The rough shape of the eigenvalue histogram is quite similar across the four panels of Fig. 4.1, which means that empirical data contains a large amount of noise, which can be described reasonably well by any of the three null models. Interesting discrepancies are present in terms of the leading eigenvalue: the empirical value is very similar to the shuffled and r-random values, while higher than the u-random value. This shows that the overall tendency towards similarity is smaller in the uniformly-random case than in the other three cases. This is understandable given that shuffling and restricted randomness reproduce the feature-level non-uniformities, which in turn are responsible for an overall level of similarity which is higher than what is expected from uniform randomness [8] (Chap. 2), leading to an enhanced global mode.

Very important are the empirical outliers in Fig 4.1(a), which encode empirical structure that is independent of feature-level nonuniformities. The two higher outliers are larger than the bulk boundary as predicted by any of the three null models, while the other two appear compatible with the random bulk predicted by uniform randomness. This highlights the importance of choosing the appropriate null model, since this determines the position of the boundary between noise modes and structural modes along the λ axis, which in turn decides how many empirical eigenmodes are to be regarded as structurally relevant on the higher λ side of this boundary. It appears that the position of this boundary is somewhat different for the three null models, but this is hard to evaluate only based on Fig. 4.1, due to limitations inherent in the binning.

Fig. 4.2(a) overcomes these limitations by showing the subleading eigenvalue distribution for the three null models, in parallel with the leading eigenvalue distributions in Fig. 4.2(b), where the colors associated to the three null models are the same as those in Fig. 4.1. For comparison, the empirical eigenvalues are shown by the vertical (red) lines in the upper bands of Fig. 4.2. Each λ_1 and λ_2 distribution is produced numerically by sampling $n = 1000$ sets of cultural vectors from the statistical ensemble of the respective null model. It appears that shuffling and r-random show essentially the same λ_2 distribution, while for u-random this is located at higher values. Since λ_2 sets the boundary for the random bulk, more empirical eigenmodes are to be regarded as structurally relevant with respect to a null model based on shuffling or restricted randomness, rather than on uniform randomness. Choosing between shuffling and r-random appears appropriate, since they are consistent with empirical data in terms of the leading eigenvalue, as noted before, now confirmed in a more statistically reliable way by Fig. 4.2(b). Such a choice is compatible with the idea of focusing on the empirical structure that is present independently of feature-level non-uniformities, which are expected to

strongly depend on how the associated questions and the possible answers are formulated and much less on authentic properties of the real social system from which the data is extracted. With respect to either the shuffled or the r-random λ_2 distribution, all four empirical outliers noted in Fig. 4.1(a) appear statistically significant, with a departure of at least two standard deviations from the mean.

On the other hand, based on Fig. 4.2(b), the empirical leading eigenvalue also appears statistically compatible with both shuffling and restricted randomness, but closer to the mean of the former. This, however, deserves a closer inspection, due to the limitations inherent in the binning of Fig. 4.2(b). Fig. 4.3 focuses on the shuffled and r-random λ_1 distributions, giving a better impression of how well either null model predicts the empirical leading eigenvalue based on partial information about trait frequencies. It appears that, due to the sharpness of the shuffled λ_1 distribution, the empirical value is actually not statistically compatible with it, while it is clearly compatible with the r-random distribution. For this reason, we choose restricted randomness as the appropriate null model. Note that, for visual purposes, the bins are chosen to be much smaller for the shuffled than for the r-random distribution – both histograms contain $n = 1000$ entries, one for each random matrix sampled from the respective ensemble.

Finally, it is worth repeating the analysis on empirical cultural states constructed from two more datasets, namely the General Social Survey [22] (GSS) – Fig. 4.4(a) – and Jester [23] (JS) – Fig. 4.4(b). Both datasets are also formatted according to the procedure described in Ref. [7] (Chap. 1), leading to $F = 122$ features for GSS and to $F = 128$ features for JS. The two figures follow the format of Fig. 4.2(a), since this emphasizes the empirical outliers and their departure from the subleading eigenvalue distributions of the three null models – although, at this point, the choice has already been made in favor of restricted randomness, the other two distributions are also shown for consistency. Both the GSS and JS eigenvalue spectra show outliers that are significantly larger than what is expected based on the r-random null model: three such outliers are present for GSS and four for JS. The deviating eigenvalues are, on average, larger for JS than for EBM, and higher for EBM than for GSS. – note that the axis ranges of Figs. 4.4(a), 4.4(b) and 4.2(a) are not the same.

Based on the results above, one can say that the empirical structure captured by matrices of cultural similarity is generally recognizable via eigenvalues that are significantly larger than what is expected based on a null hypothesis accounting for empirical trait frequencies: they are significantly higher than the subleading eigenvalue and much lower than the leading eigenvalue expected from this null hypothesis. For the rest of this study, the eigenpairs (eigenvector-value pairs) associated to these deviating eigenvalues will often be referred to as “structural modes”.

4.3 Two interpretations of structural modes

This section explores possible ways of interpreting the structural modes of culture described above. To begin with, certain aspects of linear algebra are emphasized, in relation to the diagonalization of similarity matrices, which justify an interpretation of structural modes as group modes, like in the context of correlation matrices. Then, two hypotheses are formulated: first, that structural modes are just the effect of correlations between cultural features, thus only retaining information about how the associated questions/items are chosen; second, that structural modes are an effect of genuine groups or grouping tendencies among the individuals, thus retaining information about the social system from which the data is extracted. This leads to probabilistic formulations of the two hypotheses in a very simplistic setting: the correlations-only scenario is realized as the “fully-connected Ising” (FCI) model in Sec. 4.3.1, while the groups scenario is realized as the “symmetric two-groups” (S2G) model in Sec. 4.3.2. Finally, in Sec. 4.3.3, the mathematical properties of the two models are studied in order to check that they behave as expected and to better emphasize their differences.

It is instructive to first consider some elementary, but important mathematical properties of the eigenvalues λ_l and the associated eigenvectors v_l satisfying Eq. (4.2), since they provide important hints towards how the structural modes are to be interpreted. For the sake of clarity, the following explanations make use of the term “individual” as a replacement for “cultural vector”, although most of the concepts presented are also valid, at least mathematically, for similarity matrices constructed from randomly generated cultural vectors, based on any probabilistic model.

Since the eigenvectors v_l have only real entries and form an orthonormal basis, one can write any real vector w with N entries as a linear combinations of the eigenvectors:

$$w = \sum_{l=1}^N \alpha_l v_l, \quad (4.3)$$

with real coefficients α_l . The rest of this argument is restricted to unit vectors w , which satisfy $\sum_{i=1}^N w_i^2 = 1$, which can be translated as $\sum_{l=1}^N \alpha_l^2 = 1$ in terms of the eigenvectors’ coefficients. This encompasses all the eigenvectors $w = v_l, \forall l$ as special cases. Moreover, let us define the following scalar quantity:

$$S = \sum_{i=1}^N \sum_{j=1}^N w_i s_{ij} w_j, \quad (4.4)$$

as the double contraction of the similarity matrix s with the vector w . By means of Eq. 4.4 and Eq. 4.3, for any vector w (including the special cases when this entirely matches one of the eigenvectors v_l) every entry of w becomes associated to one of the individuals based on which the similarity matrix s is computed. Thus,

w can be seen as a (normalized) linear combination of the N individuals. S can be then interpreted as the self-similarity of any normalized linear combination w , since every pairwise similarity s_{ij} is multiplied by the numbers w_i and w_j attached to individuals i and j . For any normalized w , one can show that:

$$S = 1 + 2 \sum_{i=1}^{N-1} \sum_{j=1+1}^N w_i s_{ij} w_j, \quad (4.5)$$

which immediately follows from the fact that $s_{ii} = 1, \forall i$, which is a direct consequence of how the similarity is defined in Eq. (4.1). Note that $S = 1$ whenever w gives a strength of 1 to one individual and 0 to all the other, which supports the interpretation of S as a self similarity. It is also important to note, from Eq. (4.5), that S is larger when w is such that pairs of entries (i, j) with the same sign correspond to higher values of s_{ij} and higher values of $|w_i w_j|$, while pairs with opposite signs correspond to lower values of s_{ij} and lower values of $|w_i w_j|$.

The largest self-similarity S is attained when the linear combination w , among all unit vectors, takes the form of the eigenvector v_1 with the largest associated eigenvalue λ_1 , corresponding to $\alpha_l = \delta_l^1, \forall l$. This largest self-similarity value is actually equal to the largest eigenvalue: $S = \lambda_1$. This is shown by plugging Eq. (4.2) and Eq. (4.3) into (4.4) and using the normalization condition, leading to:

$$S = \sum_{l=1}^N \alpha_l^2 \lambda_l. \quad (4.6)$$

More generally, one can see here that each eigenvector v_l with the l th highest eigenvalue λ_l , corresponding to $\alpha_{l'} = \delta_{l'}^l, \forall l'$, is such that it gives the largest possible value of $S = \lambda_l$, while also being normalized and orthogonal to all eigenvectors $v_{l'}$ with $\lambda_{l'} > \lambda_l$. When confronting this with the insights provided by Eq. (4.5), one realizes that any subset of individuals with strong, internal similarities is captured by one of the eigenmodes, whose eigenvalue is larger if the overall level of internal similarity is higher. Moreover, the eigenvector entries of these strongly similar elements will have the same sign and the highest absolute values.

By combining the above with the findings of Sec. 4.2, a more complete interpretation is obtained for structural modes: they are the normalized linear combinations of the individuals, orthogonal to each other and to the global mode, with the highest possible self-similarities, of which with the lowest is significantly higher than what is expected from restricted randomness. Each of these structure modes could indicate the presence of a group of highly similar individuals, which is why in the context of time-series analysis they are often called “group modes” [14]. Although it is not clear how a linear combination of individuals (or of cultural vectors) should be expressed in terms of cultural traits and features, this is not important for this study and does not affect the above arguments.

An alternative interpretation of structural modes comes from realizing that social surveys are imperfect, in the sense that one cannot guarantee the absence of overlaps or of similarities between the variables that are used. These translate to correlations between cultural features, which have been noticed in previous studies [5, 6, 7] (Chap. 1) and which are specific to the design of each dataset. It is conceivable that feature correlations, if strong enough, could induce artifactual structural modes themselves. For example, if a large fraction of the associated items or questions are designed such that they are mostly sensitive to the same underlying degree of freedom, the similarity between individuals responding to any of these items in a certain way will be high, since these individuals will likely respond to all the other similar items in the same way. It appears likely that this behavior would be captured by a structural mode. If this is the mechanism behind the structural modes shown in Sec. 4.2, it means that they do not provide information about the inherent organization of real-world culture, but just about the design of the “instrument” used to “measure” culture. On the other hand, to make things more complicated, feature-feature correlations may also be a consequence of group structure.

It is thus crucial to understand the extent to which structural modes of culture are due to the details of the experimental setting and to what extent they are due to authentic groups that are recognizable in the real world regardless of such details. This study makes a first step in this direction, by translating the two scenarios as mathematical, probabilistic models capable of generating (sets of) cultural vectors that are governed either by a coupling between cultural features (Sec. 4.3.1) or by a grouping tendency (Sec. 4.3.2). These models are designed to work without any empirical input, in the simplest conceivable setting, consisting of F binary features – it does not matter whether these features are regarded as ordinal or nominal, since the two types of similarity contributions are equivalent if there only $q = 2$ traits available, as can be seen from Eq. (4.1). For each feature, the two traits are marked as -1 and $+1$ – although the former should be mapped to 0 when computing similarities between vectors, if features are to be regarded as ordinal. Each of the two models defines a statistical ensemble and an associated cultural space distribution – in the language of Refs. [7, 8] (Chaps. 1 and 2) – according to which cultural vectors can be drawn in random, but non-uniform way. Both statistical ensembles are defined such that each feature-level probability distribution is uniform – the two traits have an equal probability of 0.5 attached. Note that, although both models are probabilistic in nature, neither of them is intended as a null model, since neither makes use of information from empirical data nor is it intended for direct, quantitative comparisons to empirical data, nor to be realistic to any extent. They are toy-models, intended to prove certain principles and provide certain insights about correlations and groups in the context of cultural states. Nonetheless, they do provide an arena for studying and developing certain mathematical tools in a highly controlled setting, tools that can be later used for studying empirical data.

4.3.1 The feature-feature correlations scenario

This section explains the “fully-connected Ising” (FCI) model, in the context of generating (sets of) cultural vectors in a stochastic way. The purpose of this probabilistic model is to enforce a certain level of correlation across all pairs of cultural features, controllable via one parameter, but as little as possible in addition. This can be done by properly choosing the probability distribution p taking as support the set of possible cultural vectors with F binary features, or, in other words, the set of possible spin configurations $\vec{S}$ with F lattice sites. Note that the support of this distribution has 2^F elements, which is the number of sites/points of the “cultural space”, according to the formalism in Ref. [7] (Chap. 1).

One needs to choose the maximally-random (thus minimally biased) probability distribution p that entails a certain level of feature-feature correlations. This is found by maximizing the Shannon entropy (Eq. (4.15)) subject to two constraints: one enforcing the normalization of the probability distribution (Eq. (4.16)), the other enforcing the overall level of pairwise coupling between cultural features (Eq. (4.17)). This procedure is a realization of maximum-entropy inference introduced in Ref. [18], and is described in detail in Sec. 4.A. The resulting probability distribution can be expressed as:

$$p(\mu, F, F_+) = \frac{1}{Z(\mu)} \frac{F!}{F_+!(F - F_+)!} \exp \left[\frac{\mu}{2} ((2F_+ - F)^2 - F) \right]. \quad (4.7)$$

This gives the (total) probability attached to all cultural vectors with F_+ out of F traits marked as “+” or “+1”, where μ is the parameter controlling the overall level of coupling between features. Moreover, $Z(\mu)$ is a normalization factor, namely the partition function in Eq. (4.22). Note that $\sum_{F_+=0}^F p(\mu, F, F_+) = 1.0$, since the expression combines the probability of different possible configurations with the same F_+ , which, due to symmetry reasons are equally likely. There are $F!/(F_+!(F - F_+)!)$ such configurations (the “density of states”) for each F_+ .

The model is mathematically equivalent to the Ising model of magnetism on a fully connected lattice [19], described in the canonical ensemble, with the parameter μ replacing the ratio between spin-spin coupling and temperature, which controls for the overall level of alignment between spins. This parallel does not come as a surprise: for any statistical physics ensemble defined by the averages of certain, externally controlled/measured (physical) quantities, the mathematical derivation can be formulated in terms of maximum-entropy inference [18], which ultimately provides a statistical, information-theoretic justification of minimum-bias as a replacement for assumptions like “ergodicity”. Due to this parallel, the nomenclature related to spins is sometimes used instead of that related to cultural features.

Based on Eq. (4.7), one can derive the expression for the correlation between

any two features:

$$C(\mu, F) = \frac{1}{Z(\mu)} \sum_{F_+=0}^F \frac{(F-2)! ((2F_+ - F)^2 - F)}{F_+!(F - F_+)!} \cdot \exp \left[\frac{\mu}{2} ((2F_+ - F)^2 - F) \right], \quad (4.8)$$

based on the entire statistical ensemble. The details of this derivations are also given in Sec. 4.A.

In Eq. (4.7) and Eq. (4.8), the coupling parameter is positive: $\mu \in [0, \infty)$. Physically, this corresponds to ferromagnetism, meaning that alignment between spins is favored, a tendency which is enhanced with increasing μ . Using Eq. (4.7), one can check that, for vanishing coupling $\mu = 0$, the probability of choosing a configuration with a given F_+ is directly proportional to the number of such configurations, which is specified by the binomial coefficient preceding the exponential. As μ is increased, more emphasis is given to configurations with unequal numbers of -1 and $+1$ traits, at the expense of configurations that are more balanced. Using Eq. (4.8), one can also check that the correlation $C(\mu, F)$ increases with increasing coupling μ , as expected, and that $C(0, F) = 0$ for any F .

4.3.2 The group structure scenario

This section explains the “symmetric two-groups” (S2G) model, in the context of generating (sets of) cultural vectors in a stochastic way. This probabilistic model enforces an organization of cultural vectors in terms of two, equally sized groups, with high similarities within groups and low similarities between groups. The model defines a probability distribution p taking as support the same set of possible cultural vectors as in Sec. 4.3.1: the cultural space defined by F binary features, with 2^F configuration. One of the groups is “centered” around the configuration with a -1 trait with respect to each feature, while the other group is centered around the opposite configuration, having a $+1$ trait with respect to each feature. The model is designed such that all features contribute equally to the group structure. As a consequence, this induces a certain level of correlation over all pairs of cultural features. The strength of these correlations is controlled by the same free parameter that controls the strength of the group structure.

According to the S2G model, every cultural vector that is generated is first randomly assigned to one of the two groups, with equal probabilities. These two groups are denoted as the “ -1 ” group and the “ $+1$ ” group. Then, at the level of every feature, the trait is randomly and independently chosen among the two possibilities, but with unequal probabilities: the trait with the same sign as the group is chosen with probability $1 - 2\nu$, while the trait with the opposite sign is chosen with probability 2ν . Here, $\nu \in [0, 0.25]$ is the free model parameter controlling the strength of the group structure: lower ν values imply stronger group structure and stronger correlations between features, as made more explicit

by Eq. (4.10). From this procedure, it follows that, at the level of every feature, each generated trait falls under one of the following situations:

- with probability $0.5 - \nu$, it is attached to a vector belonging to group -1 and has a value of -1 ;
- with probability ν , it is attached to a vector belonging to group -1 and has a value of $+1$;
- with probability $0.5 - \nu$, it is attached to a vector belonging to group $+1$ and has a value of $+1$;
- with probability ν , it is attached to a vector belonging to group $+1$ and has a value of -1 .

Note that the probabilities of the four cases add up to 1.0, that the combined probability of either value is 0.5 and that the probability of either group is also 0.5.

For this model, the probability that a generated configuration has F_+ traits $+1$ is:

$$\begin{aligned} p(\nu, F, F_+) &= \\ &= \frac{1}{2} \frac{F!}{F_+!(F - F_+)!} (2\nu)^{F_+} (1 - 2\nu)^{F - F_+} [(2\nu)^{F - 2F_+} + (1 - 2\nu)^{F - 2F_+}], \end{aligned} \quad (4.9)$$

while the correlation between any two features is:

$$C(\nu) = 1 - 8\nu + 16\nu^2. \quad (4.10)$$

The mathematical derivations of Eq. (4.9) and Eq. (4.10) are given in Sec. 4.B. Note that the correlation in Eq. (4.10) behaves as expected, namely: $C(0.0) = 1$ (when the two groups are maximally dissimilar the correlation is maximal) and $C(0.25) = 0.0$ (when the two groups are indistinguishable the correlation is zero). Finally, Eq. 4.10 can be written in the form of a quadratic equation, whose solution reads:

$$\nu(C) = \frac{1 - \sqrt{C}}{4}, \quad (4.11)$$

after having taken into account that $\nu \in [0, 0.25]$. Note that the alternative, $1 + \sqrt{C}$ solution given by the quadratic formula would be valid for the $\nu \in [0.25, 0.5]$ interval, which is not used here, since it is entirely equivalent (up to an inversion) with the $\nu \in [0, 0.25]$ interval, while being relevant only when group -1 is allowed to be biased towards $+1$ traits instead of towards -1 traits, and viceversa, which is not the case here.

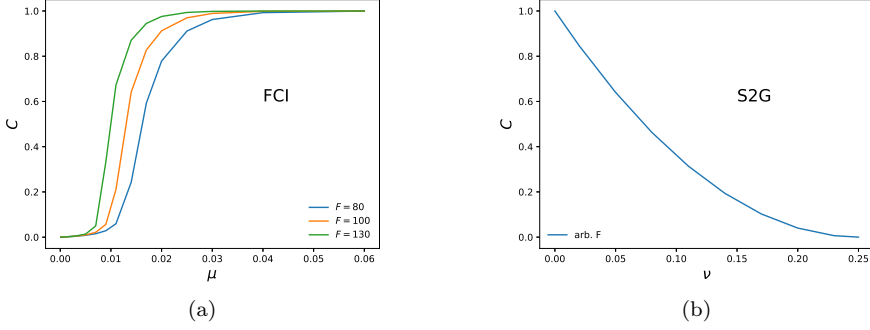


Figure 4.5: Correlation behaviour. The figure shows the dependence of the pairwise feature correlation C , first (a) on the feature-feature coupling strength parameter μ controlling the fully-connected Ising model (FCI), second (b) on the group strength parameter ν controlling the symmetric two-groups model (S2G). In the case of FCI, different curves (legend) are shown for different values of the number of features F , while in the case of S2G, a single curve is shown, which is valid for any value of F .

4.3.3 Mathematical comparison of the two scenarios

This section deals with the comparison between the FCI and the S2G models, in terms of properties that can be extracted directly from the equations in Sec. 4.3.1 and Sec. 4.3.2, without the need of randomly sampling from the two statistical ensembles. Specifically, we focus on the behavior of the feature-feature correlation (Fig. 4.5), the shape of the probability distribution (Fig. 4.6) and the symmetry breaking phase transition (Fig. 4.7) associated to each model.

Fig. 4.5 shows the behavior of the correlation between any two cultural features for the two models. First, Fig. 4.5(a) shows how the correlation entailed by the FCI model depends on the model parameter μ controlling the pairwise couplings between features, based on Eq. (4.8). Different curves correspond to different values of F . Note that the correlation increases from $C=0.0$ to $C=1.0$ as the coupling μ is increased, but it also increases as the number of features F is increased. Second, Fig. 4.5(b) shows how the correlation entailed by the S2G model depends on the model parameter ν controlling the group strength, based on Eq. (4.10). Here, the correlation decreases from $C=1.0$ to $C=0.0$ as ν is increased, which is consistent with the fact that, by construction, lower values of ν correspond to a stronger group structure. Note that the $C(\nu)$ behavior is independent of F , which is obvious from Eq. (4.10).

All the following comparisons are based on a matching of the two models in terms of the correlation level C . Specific values of μ are chosen, based on which the correlation level entailed by the FCI model $C(\mu, F)$ is computed via

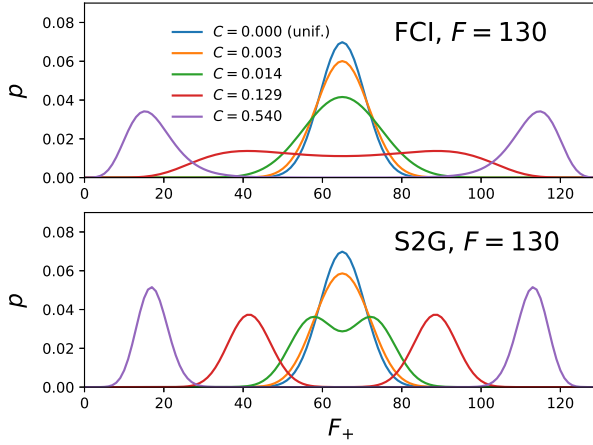


Figure 4.6: Shape of probability distribution. The figure shows the probability associated to configurations with F_+ traits +1 for $F = 130$ features, for different values of the feature-feature correlation level C (legend), for the fully-connected Ising model (FCI, top) and for the symmetric two-groups model (S2G, bottom).

Eq. (4.8), for a given F . Then, the corresponding $\nu(C)$ of S2G entailing the same correlation is calculated based on Eq. (4.11). This creates a correspondence between parameter μ of FCI and parameter ν of S2G by means of the correlation C . Since C is a number extracted from the full statistical ensemble under a specific parameterization, it can be regarded as a model parameter, namely as a replacement or remapping of μ (in the case of FCI) and of ν (in the case of S2G), which allows for a side-by-side comparison of the two models in terms of other quantities.

This μ -to- C -to- ν mapping is first exploited by Fig. 4.6, which shows the probability distributions associated to the FCI and S2G models, as described by Eq. (4.7) and Eq. (4.9) respectively. In either case, the distribution is shown for the same values of the correlation C that are listed by the legend at the top. These C values correspond to the values of the μ and ν parameters that are listed in Table 4.3.3. The calculations are based on a value of $F = 130$, which is comparable to the F values associated to the empirical cultural states used in Sec. 4.2 and Sec. 4.5.

Note that, in the limit of vanishing correlation C , the distributions of both models converge to the uniform probability distribution, which assigns to every value of F_+ a probability that is equal to the fraction of possible configurations with that many “+” traits. This uniform distribution is characterized by the existence of one maximum at the center of the F_+ axis. As the correlation C

correlation level C	FCI parameter μ	S2G paramter ν
0.000	0.000	0.250
0.003	0.002	0.237
0.014	0.005	0.221
0.129	0.008	0.160
0.540	0.010	0.066

Table 4.1: Parameter mapping. The table shows the correspondence between the correlation values C shown in Fig. 4.6, the associated μ values used for generating the FCI probability curves and the associated ν values used for generating the S2G probability curves. This correspondence is valid when $F = 130$ features are used for the FCI and S2G models.

increases, the shape of the distribution becomes wider, with two equal maxima arising on either side of the F_+ axis, whose separation also increases with increasing C . Thus, both models exhibit a symmetry breaking phase transition.

However, a close inspection of Fig. 4.6 reveals that the symmetry breaking happens later (higher values of C) for the FCI model than for the S2G model, meaning that there is a non-vanishing C interval for which FCI exhibits a unimodal behaviour, while the S2G exhibits a bimodal behaviour, interval which contains the $C = 0.014$ value. This C interval is of crucial interest for this study, since it corresponds to the correlation regime for which the symmetric group structure built into the S2G model is visible in the shape of the probability distribution, while the feature-feature coupling built into the FCI model is not strong enough to induce a qualitatively similar shape. Still, even for C values that are high enough for the FCI distribution to also show maxima, the exact shapes of the two distributions are also different, with the S2G maxima being stronger than the FCI ones (visible for $C = 0.129$ $C = 0.540$). This is a visual confirmation that the two statistical ensembles are indeed different and that the S2G ensemble has a smaller Shannon entropy than the FCI ensemble, for any, non-vanishing value of C , thus being more biased, more constrained and encoding more structure, which should manifest itself at the level of higher-order correlations (involving more than two spins/features).

A more complete picture of the phase transitions exhibited by the two models is provided by Fig. 4.7. This shows the dependence of two mathematical properties of the probability distributions in Fig. 4.6 on the model parameters. The first property, denoted here by $O_1(\gamma, F) \in [0, 1]$, is a normalized departure of either probability peak from the center of the (horizontal) F_+ axis. The second property, denoted here by $O_2(\gamma, F) \in [0, 1]$, is a normalized height of either probability peak compared to the probability at the center of the (horizontal) F_+ axis. Note that γ is a placeholder for either the μ parameter or the ν parameter, depending, respectively, on whether the FCI or the S2G model is used. Both quantities are

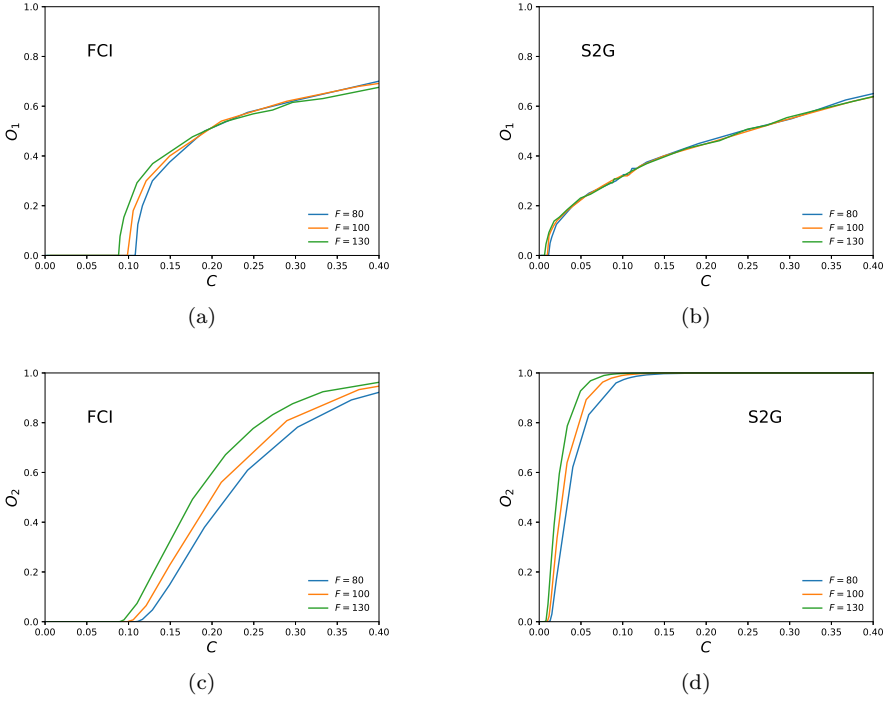


Figure 4.7: Symmetry breaking phase transitions. The figure shows the behavior of the normalized probability peak departure O_1 (top) and of the normalized peak height O_2 (bottom), as functions of the model parameters, for the fully-connected Ising (FCI, top) and the symmetric two-groups (S2G, bottom) models. Different curves corresponds to different values of the F paramter, controlling the number of features (legends). Both the μ parameter of the FCI model and the ν parameter of the S2G model are remapped to the associated correlation value C , which is shown along the horizontal axis of each plot.

zero when symmetry breaking is not present and are positive when symmetry breaking is present, giving higher values for better defined probability peaks. They can thus be used as “order parameters” characterizing the phase transition, although they are evaluated in a *a priori* way, based on the expression of the probability distribution, rather than based on configurations sampled from the associated ensemble. Mathematically, the first quantity is defined as:

$$O_1(\gamma, F) = \frac{[0.5F] - F_+^*(\gamma, F)}{[0.5F]}, \quad (4.12)$$

while the second quantity is defined as:

$$O_2(\gamma, F) = \frac{p^*(\gamma, F) - p(\gamma, F, [0.5F])}{p^*(\gamma, F)}, \quad (4.13)$$

where the square brackets stand for the “integer part” operation. Moreover, $F_+^*(\gamma, F)$ is the (integer) position along the F_+ axis of the first (lower- F_+) peak and $p^*(\gamma, F)$ is the height of this peak. At the same time, $p(\gamma, F, [0.5F])$ is evaluated according to either Eq. (4.7) or Eq. (4.9), depending on whether the quantity is evaluated for the FCI model (γ is replaced by μ) or for the S2G model (γ is replaced by ν). The value of $F_+^*(\gamma, F)$ is extracted by iteratively exploring the lower half of the F_+ axis, while evaluating $p(\gamma, F, F_+)$ according to either Eq. (4.7) or Eq. (4.9). On the other hand, $p^*(\gamma, F)$ is essentially an abbreviation for $p(\gamma, F, F_+^*(\gamma, F))$.

The four panels of Fig. 4.7 show the behaviour of O_1 for the FCI model (Fig. 4.7(a)), the behaviour of O_1 for the S2G model (Fig. 4.7(b)), the behaviour of O_2 for the FCI model (Fig. 4.7(c)) and the behaviour of O_2 for the S2G model (Fig. 4.7(d)). The dependence of either quantity on the μ parameter (for FCI) and on the ν parameter (for S2G) is translated in terms of the corresponding correlation value C , via Eq. (4.8) and Eq. (4.10) respectively. Note that the two quantities agree in terms of the correlation value for which the transition occurs, for both the FCI (Fig. 4.7(a) vs Fig. 4.7(c)) and the S2G (Fig. 4.7(b) vs Fig. 4.7(d)), for any number of features F . It is clear that the transition point comes closer to $C = 0.0$ with increasing F for both models. Finally, Fig. 4.7 shows that, independently of F , the transition point of S2G is located at lower values of C than that of FCI.

4.4 Discriminating between the two interpretations

This section investigates the signatures of the two structural scenarios introduced in Sec. 4.3, from a spectral analysis and random matrix perspective, with the purpose of identifying quantities that can differentiate between the two underlying hypotheses: feature-feature correlations vs group structure. To this end, sets of

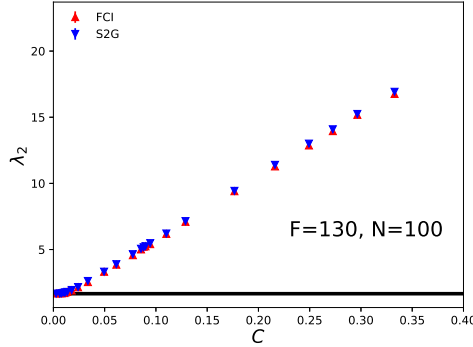


Figure 4.8: Behaviour of subleading eigenvalue (λ_2). The figure shows how λ_2 depends on the correlation level C for the fully-connected Ising (FCI, red, upward triangles) and for the symmetric two-groups (S2G, blue, downward triangles) models. For each C value, for each of the two models, an averaging is performed over 80 sets of cultural vectors independently sampled from the respective ensemble – the vertical bar associated to each point shows the interval spanned by one standard mean error on each side of the mean. The black, horizontal lines show, for comparison, the mean λ_2 expected based on uniform randomness, along with the width of the λ_2 distribution – one standard deviation on each side – where the calculations are based on 60 sets of cultural vectors generated via uniform randomness – these lines do not imply that, for uniform randomness, the correlation C (which actually vanishes by construction) can be arbitrarily large.

cultural vectors are numerically sampled from the two ensembles and similarity matrices are computed, based on Eq. (4.1). Since both the FCI and S2G ensembles are such that the (marginal) feature-level probability distributions are uniform, restricted randomness (see Sec. 4.2) is equivalent to uniform randomness as a null model (at least if the number of cultural vectors N is reasonably high) with respect to which structure is to be evaluated. Thus, for simplicity, uniform randomness (u-random) is used as a null model in this section. All comparisons made here make use of matching the feature-feature coupling parameter μ of FCI and the group strength parameter ν of S2G in terms of the correlation level C , as described in Sec. 4.3.3. Moreover, the number of features and the number of cultural vectors are $F = 130$ and $N = 100$ for all the FCI, S2G and u-random cultural states generated and used for the figures of this section.

The most obvious quantity that could conceivably discriminate between the FCI and the S2G models is the subleading eigenvalue λ_2 , or the extent to which this goes above the uncertainty range predicted by uniform randomness. Fig. 4.8 shows the dependence of λ_2 on the correlation level C for FCI (red) and S2G

(blue), while the horizontal black lines show the u-random uncertainty range (the mean value and 1 standard deviation on each side of the mean), as a compact replacement of the distributions shown in Fig. 4.2(a), Fig. 4.4(a) and Fig. 4.4(b) – as mentioned in the figure caption, these lines are not meant to give any information about the correlation level of the u-random null model, nor about realized correlations based on specific sets of vectors sampled from the ensemble. Surprisingly, λ_2 does not distinguish between the FCI and the S2G models, for any given correlation level C , since the average λ_2 values clearly overlap. At the same time, λ_2 (for both models) does depart significantly from the null model expectations. This explicitly shows that empirical structural modes such as those identified in Sec. 4.2 can actually be triggered by feature-feature correlations alone, at least in certain cases (those for which the simplistic setting behind the FCI and S2G models is reasonably representative). Thus, empirical eigenvalues that significantly depart from what is expected based on the null hypothesis do not automatically indicate groups. In the light of Sec. 4.3, Fig. 4.8 also implies that the subleading eigenmodes of matrices produced via FCI are associated, on average, to the same self-similarity as those of matrices produced via S2G, for a given correlation level. This appears counter-intuitive, since the low- C presence of symmetry breaking for S2G makes it much easier to identify two, well separated groups, one for each side of the F_+ axis of Fig. 4.6. However, a closer inspection of the probability distributions in Fig. 4.6 reveals that FCI is more likely to produce, even in the absence of symmetry breaking, cultural vectors that are at one extreme or the other (almost fully populated with $+1$ traits or with -1 traits). These extremal configurations are much more representative, or “central”, for the configurations that are possible on the respective side of the F_+ axis. Also note that the values of C used in Fig. 4.8 are the same for FCI and S2G and the same as those used in Fig. 4.9 and Fig. 4.10 described below. For each FCI and S2G point in any of these plots, explicit averaging over the sampled sets of cultural vectors is only performed with respect to the quantity associated to the vertical axes. For the correlation level C , associated to the horizontal axes, we simply use the analytically-computed, ensemble-level value, for the given parameterization of the model (Eq. (4.8) and Eq. (4.10)).

Sec. 4.C shows, in a manner similar to Fig. 4.8, the behavior of the largest and third largest eigenvalues – λ_1 and λ_3 respectively – for the FCI and S2G models, in comparison to the u-random null model. The analysis there makes it clear that the λ_1 and λ_3 are both compatible with the null hypothesis. Thus, all or most of the structural information of cultural states generated from either the FCI or the S2G model is captured by the (λ_2, v_2) eigenpair. Since λ_2 cannot discriminate between the two scenarios, this means that all or most discriminating power is encoded in the associated eigenvector v_2 , which is the focus of the rest of this section.

Based on Sec. 4.3.3 and in particular on Fig. 4.6, one can say that, for the interesting correlation interval where FCI does not exhibit symmetry breaking while S2G does, configurations that are on one side of the F_+ axis and are gen-

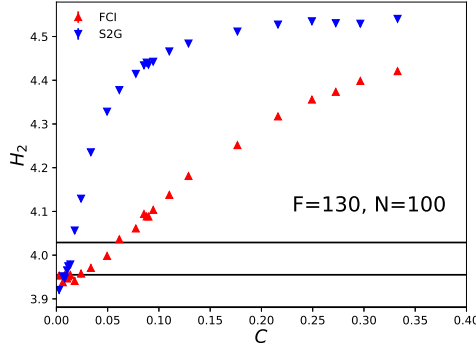


Figure 4.9: Behaviour of uniformity H_2 associated to subleading eigenvalue. The figure shows how H_2 depends on the correlation level C for the fully-connected Ising (FCI, red) and for the symmetric two-groups (S2G, blue) models. For each C value, for each of the two models, an averaging is performed over 80 sets of cultural vectors independently sampled from the respective ensemble – the vertical bar associated to each point shows the interval spanned by one standard mean error on each side of the mean. The black, horizontal lines show, for comparison, the mean H_2 expected based on uniform randomness, along with the width of the H_2 distribution – one standard deviation on each side – where the calculations are based on 60 sets of cultural vectors generated via uniform randomness – these lines do not imply that, for uniform randomness, the correlation C (which actually vanishes by construction) can be arbitrarily large.

erated with S2G have a much more equal fraction of traits of a certain sign than those generated with FCI. These S2G configurations should thus have a much more equal contribution to the structural mode (λ_2, v_2) than FCI configurations, so the associated v_2 entries should be much more equal for S2G than for FCI. Given the symmetric nature of both models, it follows that the absolute values of all the v_2 entries should be much more equal for S2G cultural states than for FCI ones, while, in either case, the entries associated to cultural vectors on different sides of the F_+ axis would (typically) have different signs. This reasoning suggests that the difference between FCI and S2G would be captured by a quantity that evaluates the overall extent of “equality” of the absolute values of the entries of the v_2 eigenvector, or, in other words, the eigenvector “uniformity”. Since these entries are normalized via $\sum_{i=1}^N |v_i|^2 = 1$ for any eigenvector v_l , the Shannon entropy is a natural quantity for evaluating the uniformity. This leads to the definition of “eigenvector entropy” H_l associated to the l th highest eigenvalue

λ_l , as a measure of uniformity:

$$H_l = - \sum_{i=1}^N |v_l^i|^2 \log |v_l^i|^2 \quad (4.14)$$

where v_l^i is the i th entry of the eigenvector associated to λ_l – note that this quantity was also used in Ref. [17], which cites Ref. [24].

Fig. 4.9 shows the behavior of the eigenvector entropy H_2 associated to the second highest eigenvalue λ_2 , in a format very similar to that of Fig. 4.8. This confirms that H_2 discriminates well between the two models, with S2G showing clearly higher H_2 values than FCI as long the correlation level does not come arbitrarily close to $C = 0.0$. Moreover, comparing the two profiles with the u-random one- σ band reveals that the structure of S2G becomes incompatible with the null-hypothesis for much lower correlation values than the structure of FCI. However, for either model, the $H_2(C)$ curve does not show the sudden increase that one would expect based on the phase transitions described in Sec. 4.3.3, in the manner they are exhibited by the more theoretical $O_1(C)$ and $O_2(C)$ curves in Fig. 4.7.

The smoothness of the $H_2(C)$ curves is actually related to the fact that, for the low- C regime, where λ_2 is highly compatible with the null hypothesis, H_2 is typically not the second highest eigenvector entropy, although it is associated to the second highest eigenvalue. This suggests a definition of H'_l as the l th highest eigenvector entropy, independently of the associated eigenvalue. Fig. 4.10 is a modification of Fig. 4.9, with H'_2 used as a replacement for H_2 for the vertical axis, affecting all the FCI, S2G and u-random calculations. Note that, unlike in Fig. 4.9, the sudden changes in Fig. 4.7 are now reflected in Fig. 4.10. Moreover, the transition points at $F = 130$ in Fig. 4.7 seem to be well reproduced in Fig. 4.10, while the FCI and S2G shapes of the $H'_2(C)$ curves are quite similar to those of $O_2(C)$, which are related to the height of the probability distribution peaks. Finally for higher C values, each $H'_2(C)$ curve in Fig. 4.10 is almost identical to the associated $H_2(C)$ in Fig. 4.9, so strong structure makes it very likely that the eigenvector of the second highest eigenvalue has the second highest entropy, and H'_2 is effectively equivalent to H_2 .

The considerations above strongly suggest that a significant departure of the eigenvector entropy from the null model expectation is a good indication that the eigenvector encodes information about a group or a grouping tendency. In the simplistic (binary, marginally-uniform) setting of the FCI and S2G models, one could define the presence of groups in a theoretical, a priori way via the presence of maxima (and symmetry breaking) in the probability distribution over the F_+ axis: when maxima are present, most vectors sampled from the distribution can be unambiguously recognized as belonging to one of the two groups, based on their F_+ value. Under this interpretation, within the interesting C interval for which S2G exhibits groups and FCI does not, the eigenvector entropy and its departure from randomness expectations is crucial for deciding, in a a-posteriori

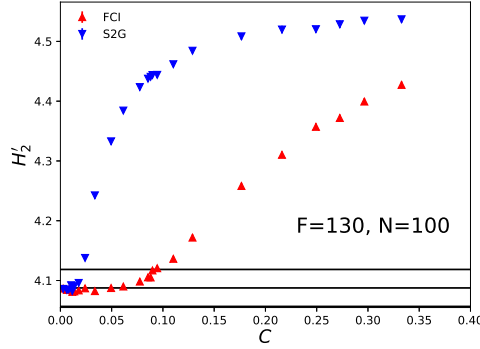


Figure 4.10: Behavior of subleading uniformity H'_2 . The figure shows how H'_2 depends on the correlation level C for the fully-connected Ising (FCI, red) and for the symmetric two-groups (S2G, blue) models. For each C value, for each of the two models, an averaging is performed over 80 sets of cultural vectors independently sampled from the respective ensemble – the vertical bar associated to each point shows the interval spanned by one standard mean error on each side of the mean. The black, horizontal lines show, for comparison, the mean H'_2 expected based on uniform randomness, along with the width of the H'_2 distribution – one standard deviation on each side – where the calculations are based on 60 sets of cultural vectors generated via uniform randomness – these lines do not imply that, for uniform randomness, the correlation C (which actually vanishes by construction) can be arbitrarily large.

way, whether groups are present or not.

4.5 Revisiting the empirical data

The findings of Sec. 4.4 point out the importance of the eigenvector entropy, in addition to the eigenvalue, for deciding whether a structural mode qualifies as an authentic group mode or not. Thus, the two quantities should be used together for a second, more detailed inspection of the empirical data in Sec. 4.2. This is the purpose of the current section. The empirical similarity matrices are computed based on the same three sets of $N = 100$ cultural vectors used in Sec. 4.2, constructed from Eurobarometer (EBM), General Social Survey (GSS) and Jester (JS) data.

Fig. 4.11 shows a scatter of the empirical eigenpairs of the EBM matrix, where the horizontal axis is associated to the eigenvalue λ , while the vertical axis is associated to the eigenvector entropy H . The global mode eigenpair is highlighted

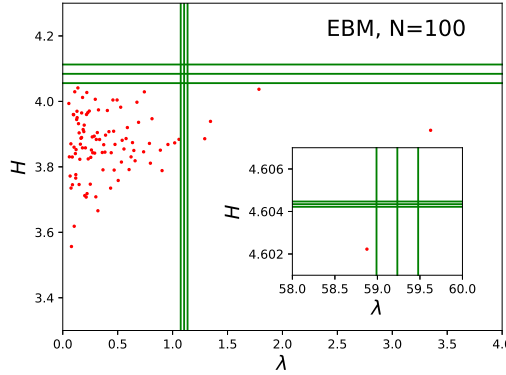


Figure 4.11: Eigenvalues and eigenvector entropies for empirical data. Every point corresponds to an empirical eigenpair, with the eigenvalue λ shown along the horizontal axis and the eigenvector entropy H shown along the vertical axis. The inset focuses on the leading eigenvalue, which also corresponds to the highest eigenvector entropy. The vertical lines in the main plot and in the inset show, respectively, the widths (one-standard deviation on each side of the mean) of the subleading and leading eigenvalue distributions, based on restricted randomness. The horizontal lines in the main plot and in the inset show, respectively, the widths (one-standard deviation on each side of the mean) of the second highest and highest eigenvector entropy distributions, based on restricted randomness. The vertical lines are not intended to provide any information about the eigenvector entropies associated to the respective eigenvalues, while the horizontal lines are not intended to provide any information about the eigenvalues associated to the respective eigenvector entropies. The figure is based on the same, Eurobarometer (EBM) data with $N = 100$ cultural vectors used in Figs. 4.1, 4.2 and 4.3.

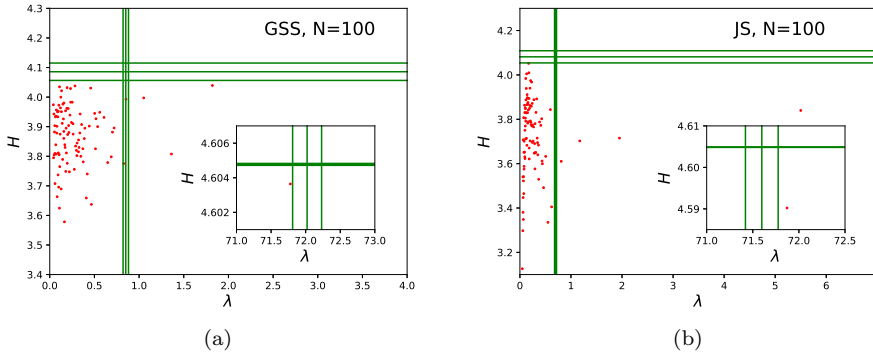


Figure 4.12: Eigenvalues and eigenvector entropies in other empirical datasets. Fig. (a) is based on the General Social Survey (GSS) data with $N = 100$ cultural vectors used for Fig. 4.12(a), while Fig. (b) is based on the Jester (JS) data with $N = 100$ cultural vectors used in Figs. 4.12(b). In each plot, every point corresponds to an empirical eigenpair, with the eigenvalue λ shown along the horizontal axis and the eigenvector entropy H shown along the vertical axis. Each inset focuses on the leading eigenvalue, which also corresponds to the highest eigenvector entropy. The vertical lines in each main plot and in each inset show, respectively, the widths (one-standard deviation on each side of the mean) of the subleading and leading eigenvalue distributions, based on restricted randomness. The horizontal lines in each main plot and in each inset show, respectively, the widths (one-standard deviation on each side of the mean) of the second highest and highest eigenvector entropy distributions, based on restricted randomness. The vertical lines are not intended to provide any information about the eigenvector entropies associated to the respective eigenvalues, while the horizontal lines are not intended to provide any information about the eigenvalues associated to the respective eigenvector entropies.

by the inset. In the main plot, the vertical lines show the average and the $1\text{-}\sigma$ band of what one may expect for the subleading eigenvalue λ_2 , based on the r -random null model, which reproduces, on average, the empirical trait frequencies (see Sec. 4.2). In the inset, the vertical lines show the same type of information for the leading eigenvalue λ_1 . The horizontal lines in the main plot and the inset show the average and the $1\text{-}\sigma$ band of what one may expect for, respectively, the subleading entropy H'_2 and the leading entropy H'_1 , based on the r -random null model. Note that, as anticipated in Sec. 4.4, the subleading entropy is usually not associated to the subleading eigenvalue, while the leading entropy appears to always be associated to the leading eigenvalue.

The four structural modes identified based on Fig. 4.2(a) are also visible in the

main plot of Fig. 4.11, to the right of the vertical r-random band. Importantly, all their eigenvector entropies are below the horizontal r-random band, suggesting that neither of them qualifies as a group mode. Actually, all the bulk EBM eigenpairs are also below the r-random band, and thus compatible with the null hypothesis in terms of the uniformity of eigenvector entries. Also note that the leading eigenvector entropy is significantly smaller than what the null model predicts, but this difference is much smaller than the difference between the leading eigenvector entropy and the subleading one. This means that the contributions of different cultural vectors to the global mode are less equal than expected based on randomness, but much more equal than the contributions to any of the structural modes.

The analysis in Fig. 4.11 is also applied to the other datasets and the results are presented in Fig. 4.12, with Fig. 4.12(a) showing the results for GSS data and Fig. 4.12(b) showing the results for JS data. In both cases, the results are similar to those of EBM data: the structural modes do not show a higher eigenvector uniformity than what is expected based on the null model, nor do any of the smaller- λ modes, while the eigenvector uniformity of the global mode is smaller than what is expected based on the null model, but much higher than what is expected or realized for the structural modes and the random modes. In the light of Sec. 4.3 and Sec. 4.4, these results suggests that structural modes of empirical matrices of cultural similarity are not due to authentic group structure, but to correlations between cultural features originating in arbitrary similarities between the questions or items composing the dataset. However, as discussed in Sec. 4.6, such a conclusion would be premature, implicitly relying on assumptions about cultural groups that might be too strong.

4.6 Discussion

This was the first study where empirical matrices of cultural similarity between individuals were analyzed from a random matrix perspective, allowing for a separation of structurally irrelevant eigenmodes from the structurally relevant ones. The statistical significance of the latter, here referred to as “structural modes”, was demonstrated in Sec. 4.2, using a detailed numerical approach of explicitly sampling configurations from three null models. Among these three, the “restricted randomness” model, first proposed here, was concluded to be the most appropriate for later use. Restricted randomness enforces, in a flexible way, the non-uniformity inherent in each cultural feature, as this is assumed to be mostly a consequence of experimental design rather than a consequence of system-specific structure. As a consequence, this null model reproduces well the leading eigenvalue of the empirical matrix, which is interpreted as the “global mode”. By using this null model, meaningful empirical structure is implicitly defined via the inhomogeneities present in the cultural space distribution [7, 8] (Chaps. 1 and 2) that cannot be expressed in terms of the feature-level inhomogeneities.

A central question for the rest of the study was whether the structural modes identified in Sec. 4.2 are just signatures of correlations between cultural features or, more interestingly, signatures of cultural groups. The former hypothesis goes along with the idea that some of the items in the questionnaire are similar to each other. The latter hypothesis goes along with the idea of coexistence, within the geographical region from which the empirical data was obtained, of several types of individuals, where each type could correspond, for instance, to a certain political affiliation, assuming that each affiliation comes along with a certain set of values, opinions or beliefs. Even more interesting is the possibility that structural modes correspond to groups that form around cultural prototypes [6, 8] (Chap. 2) associated to a small number of universal “rationalities” or “ways of life” [20]. This hypothesis has been shown to be compatible with some generic structural properties of culture, provided that prototype mixing is in place [8] (Chap. 2).

We approached this question by designing, in the simplest possible setting, two probabilistic toy models that implement the “correlations” scenario and the “groups” scenario (see Sec. 4.3) named “FCI” (Sec. 4.3.1) and “S2G” (Sec. 4.3.2) respectively. These models and the associated scenarios are not mutually exclusive: the presence of groups does induce correlations, while correlations, if strong enough, can also induce an “impression” of groups. However, the FCI model is conceived such that only feature-feature couplings are enforced in a manner that does not introduce any unintended assumption, by means of a maximum-entropy approach [18]. This is meant to “simulate” an overall level of pairwise similarity between the questions of a hypothetical survey, assuming that the hypothetical system from which the answers are obtained is otherwise maximally random. Moreover, there is a non-vanishing correlation interval for which (under a certain, meaningful projection) the S2G model has a bimodal probability distribution (Sec. 4.3.3), while the FCI model has a unimodal distribution. One can say that, for this interval, the group structure of S2G is manifested, while the feature-feature couplings of FCI do not yet create the impression of groups. The boundaries of this interval are well defined, via the symmetry breaking phase transition of S2G on the low-correlation side and the one of FCI on the high-correlation side.

This correlation region is exploited (Sec. 4.4) for understanding how the presence or absence of groups becomes visible via spectral analysis. In both cases, there is one eigenvalue that becomes increasingly separated from the random bulk when increasing the level of correlations between features. However, this increasing trend is, up to statistical errors arising from finite sampling, exactly the same for the FCI and S2G models, even for the above-mentioned correlation region. This suggests that the presence of deviating eigenvalues in empirical data is not a certain signature of group structure. The difference between the two scenarios becomes visible if one calculates the uniformities of the eigenvectors by means of “eigenvector entropy” (inspired by Ref. [17], where it is called “information entropy”). There is one eigenvector uniformity that, for an increasing level of correlations, becomes increasingly separated from the random bulk. This increasing

trend is significantly different for the two models, while starting in an abrupt way and replicating well, for each model, the phase transition expected on theoretical grounds. Thus, for the interesting correlation region, S2G shows a deviating eigenvector uniformity, while FCI does not. This suggests that empirical eigenmodes corresponding to authentic groups should exhibit not only an eigenvalue that is significantly higher than the null model expectation, but also an eigenvector uniformity that is significantly higher than the null model expectation.

This motivated a more detailed investigation of empirical data in Sec. 4.5, which showed that all empirical eigenvalues that are significantly higher than what can be expected based on restricted randomness are associated to eigenvector uniformities that are not significantly higher than what can be expected based on the same null model. This suggests that empirical deviating eigenvalues are signatures of correlations and not of group structure, since such correlations are known to be present, although to different extents and differently distributed in different datasets [7] (Chap. 1). One may even be tempted to reject the “cultural prototypes” hypothesis previously used in Refs. [6, 8] (Chap. 2). However, Ref. [8] (Chap. 2) clearly showed that this hypothesis is structurally compatible with empirical data only when prototype “mixing” is enforced, which means that the cultural vectors associated to different individuals are random combinations of the prototype vectors, although each vector is most often dominated by one of the prototypes. Since the S2G model used here to simulate group structure does not incorporate mixing, it is possible that group structure resulting from a “mixed prototypes” scenario is different enough to not exhibit eigenvector uniformities which are higher than expected based on the the null hypothesis.

Actually, the implementation of the “Mixed Prototype Generation” procedure of Ref. [8] (Chap. 2) is able to generate, for many parameter choices, vectors that are arbitrarily similar to one of the prototypes, as well as vectors that are balanced combinations of the prototypes. If a modified S2G model incorporating such a mixing would be formulated, this would very likely be able to induce, in the language of Fig. 4.6, probability distributions that are wider than those of S2G, showing weaker decays when approaching the $F_+ = 0$ and the $F_+ = F$ endpoints, while still different than those of FCI, for a given correlation level. These distributions might not even show a double-peak structure, and would likely preserve their shapes in the limit of $F \rightarrow \infty$ – assuming that the $[0, F]$ interval is mapped to another interval of a constant length when F increases – while the peaks of the S2G distributions become sharper with increasing F , due to the central limit theorem. It appears likely that cultural states sampled from such “mixed-S2G” distributions would only exhibit a subleading eigenvector uniformity that significantly deviates from null model expectations for correlation levels that are higher than those required by FCI: below that level, the vectors composing each of the two “groups” would have highly different levels of “centrality” within the group, leading to non-equal entries in the eigenvector capturing most of the structure. Certainly, such a mixed-S2G would come with a rather different meaning of “groups” and of “group structure” than that implicit in S2G

and recognizable via eigenvector uniformity. One would also need to find a new, eigenvector-dependent quantity that is sensitive to this different type of group structure and that can be also used for empirical data. Such considerations are left for future research.

The fact that this study used multidimensional sociological data, while heavily relying on eigenvalue decomposition, may raise the question of how the approach here is different from traditional social science research using principal component analysis [25]. Although principal component analysis heavily relies on eigenvalue decomposition, in a social science context, the former most often implies a decomposition of the matrix of covariances or correlations between the variables, while this study focuses on the matrix of similarities between individuals. This actually makes the approach here conceptually more similar to clustering methods [26], which aim at identifying group structure, while providing an optimal clustering of the given set of individuals. However, these methods do not attempt to decompose the similarity matrix and remove the irrelevant eigenmodes. In fact, following the approach of Ref. [14], the sum of the similarity matrix contributions associated to the structural modes identified here can be interpreted as a modified modularity matrix, which could provide a new method for clustering individuals via modularity maximization. Since this automatically eliminates the noise components and the common trend encoded in the global mode, such a method should be able to disentangle clusters that are not recognized by previous approaches. However, such a method might also be sensitive to false positive cluster splittings, due to structural modes possibly being artifacts of feature-feature correlations, as shown in this study (at this point, it is not clear whether this is also a problem for the method in Ref. [14], intended for matrices of correlations between time series). These aspects are also left for future research.

4.7 Conclusion

This study examined cultural structure from a new angle, relying on certain notions of random matrix theory. This provided a filtering procedure for matrices of cultural similarity between individuals, which eliminates, in a statistically rigorous way, the structurally-irrelevant components. Much effort was dedicated to the interpretation of the remaining, structurally-relevant components. Two possible interpretations were formulated and quantitatively examined. On one hand, structural components may be a consequence of overlaps between cultural variables, mainly encoding information about the experimental setup. On the other hand, they may be a consequence of a modular organization of culture, thus encoding information about cultural groups. The analysis here favored the former scenario, but this may be a consequence of the latter scenario having been formalized in a manner that is too restrictive. More work is needed for entirely rejecting or accepting the possibility that culture has a modular structure.

Appendices

4.A The fully-connected Ising (FCI) model

This section gives the details behind the mathematical expressions in Sec. 4.3.1, which introduced the fully-connected Ising model. Deriving the probability distribution p associated to this model follows the maximum-entropy approach introduced by Ref. [18]. This crucially relies on the Shannon entropy, which is a functional of the probability distribution:

$$H[p] = - \sum_{\vec{S}} p_{\vec{S}} \log p_{\vec{S}}, \quad (4.15)$$

where $\vec{S}$ denotes a generic spin configuration with F spins on a fully-connected lattice, or a generic cultural vector with F binary cultural features whose possible traits are marked as “−1” and “+1”. The value of the functional H is maximized subject to two constraints, one related to the normalization of the probability distribution over the set of possible configurations:

$$\sum_{\vec{S}} p_{\vec{S}} = 1, \quad (4.16)$$

the other related to enforcing, on average, a certain amount K of alignment:

$$\sum_{a < b} \sum_{\vec{S}} S_a S_b p_{\vec{S}} = K, \quad (4.17)$$

namely the average number of pairs of similarly labeled traits within a given configuration $\vec{S}$, where the first summation is over all distinct pairs of distinct features (or lattice sites). The maximization is done using the Lagrange multipliers technique for Eqs. (4.15), (4.16), (4.17), which implies that one should find the extrema of the following functional:

$$L[p] = H[p] - \lambda_0 \left(\sum_{\vec{S}} p_{\vec{S}} - 1 \right) - \lambda \left(\sum_{a < b} \sum_{\vec{S}} S_a S_b p_{\vec{S}} - K \right), \quad (4.18)$$

where λ_0 and λ are free parameters associated to the two constraints. By taking partial derivatives of Eq. (4.18) with respect to each $p_{\vec{S}}$ and further manipulations, one finds the following probability distribution:

$$p_{\vec{S}} = \frac{1}{Z(-\lambda)} \exp \left[-\lambda \sum_{a < b} S_a S_b \right], \quad (4.19)$$

where $Z(-\lambda)$ is a normalization factor, known in statistical physics as the “partition function”:

$$Z(-\lambda) = \sum_{\vec{S}} \exp \left[-\lambda \sum_{a < b} S_a S_b \right], \quad (4.20)$$

where one can replace the coupling parameter $-\lambda$ with $\mu > 0$ (whose positive value favors alignment as opposed to anti-alignment, which corresponds to ferromagnetism) and re-express the sum over configurations $\vec{S}$ as a sequence of sums over the possible traits of each feature S_k , leading to:

$$Z(\mu) = \prod_{k=1}^F \left(\sum_{S_k=\pm 1} \right) \exp \left[\mu \sum_{a=1}^{F-1} \sum_{b=a+1}^F S_a S_b \right]. \quad (4.21)$$

In the exponent of this expression, there are $F(F-1)/2$ terms, out of which $F_+(F-F_+)$ are equal to -1 , while the other are equal to $+1$. Based on this, after further manipulations and after taking advantage of symmetries, the partition function can be expressed as as:

$$Z(\mu) = \sum_{F_+=0}^F \frac{F!}{F_+!(F-F_+)!} \exp \left[\frac{\mu}{2} ((2F_+ - F)^2 - F) \right], \quad (4.22)$$

where the combinatorial factor (binomial coefficient) before the exponential function counts the number of configurations with the same number F_+ of $+1$ traits (the density of states). In a way rather analogous to the partition function, the double summation in the exponent of Eq. (4.19) can also be eliminated. After multiplication with the density of states, this leads to Eq. (4.7), which gives the probability of having a configuration with F_+ spins up.

On the other hand, using Eq. (4.20), Eq. (4.17) can be written as:

$$K = -\frac{\partial(\log(Z(-\lambda)))}{\partial \lambda} = \frac{\partial(\log(Z(\mu)))}{\partial \mu}, \quad (4.23)$$

while the correlation between features/spins a and b is:

$$C_{ab} = \frac{\langle S_a S_b \rangle - \langle S_a \rangle \langle S_b \rangle}{\sqrt{\langle S_a^2 \rangle - \langle S_a \rangle^2} \sqrt{\langle S_b^2 \rangle - \langle S_b \rangle^2}}, \quad (4.24)$$

where $\langle Q \rangle = \sum_{\vec{S}} Q_{\vec{S}} p_{\vec{S}}$ is the expected value of quantity Q with respect to the statistical ensemble. However, one can easily show, using Eq. (4.22) that $\langle S_a^2 \rangle = 1$ and that $\langle S_a \rangle = \langle S_b \rangle = 0$, so $C_{ab} = \langle S_a S_b \rangle = \sum_{\vec{S}} S_a S_b p_{\vec{S}}$, which combined with Eq. (4.17) leads to $\sum_{a < b} C_{ab} = K$. But due to symmetry, the expected correlation C_{ab} is the same for all pairs (a, b) , so:

$$C_{ab} = C(\mu, F) = \frac{2}{F(F-1)} K = \frac{2}{F(F-1)} \frac{\partial(\log(Z(\mu)))}{\partial \mu}, \quad (4.25)$$

for any pair (a, b) , which can also be written in the form shown by Eq. (4.8) – Eq. (4.23) was used for the last transformation in Eq. (4.25).

One should expect that $C(0.0) = 0.0$ (null correlations for null coupling), which based on Eq. (4.8), implies that the following identity holds:

$$\sum_{F_+=0}^F \frac{(F-2)!((2F_+-F)^2-F)}{F_+!(F-F_+)!} = 0, \quad (4.26)$$

which, after substitution of F_+ with k and of F with N and some further manipulations leads to the following combinatorial identity:

$$\sum_{k=0}^N \binom{N}{k} ((2k-N)^2 - N) = 0 \quad (4.27)$$

which can be shown to hold using the expressions for the binomial expansion and for the first and second moments of a binomial distribution with the probability parameter set to 0.5.

4.B The symmetric two-groups (S2G) model

This section provides the mathematical derivations of the important mathematical formulas related to the symmetric two-group model, introduced in Sec. 4.3.2. The derivations are based on the model description there.

First, we prove Eq. (4.9). On one hand, the probability that a cultural vector meant to be part of group +1 is assigned to a configuration with F_+ traits +1 is:

$$p_+^+(\nu, F, F_+) = \frac{F!}{F_+!(F-F_+)!} (1-2\nu)^{F_+} (2\nu)^{F-F_+}, \quad (4.28)$$

which is a binomial distribution with probability $1-2\nu$ for the +1 possibility and 2ν for the -1 possibility. On the other hand, the probability that a configuration meant to be part of group -1 has F_+ traits +1 is:

$$p_+^-(\nu, F, F_+) = \frac{F!}{F_+!(F-F_+)!} (2\nu)^{F_+} (1-2\nu)^{F-F_+}, \quad (4.29)$$

which is the same binomial distribution, but with inverted probabilities. Since the two groups are by construction equally likely, the combined probability of all configurations with F_+ traits +1 is:

$$p(\nu, F, F_+) = \frac{1}{2} p_+^+(\nu, F, F_+) + \frac{1}{2} p_+^-(\nu, F, F_+). \quad (4.30)$$

Inserting Eq. (4.28) and Eq. (4.29) in Eq. (4.30) leads to Eq. (4.9).

Second, we prove Eq. (4.10). The correlation coefficient of any two features a and b is given by Eq. (4.24), which, for symmetry reasons similar to the case of the FCI model, simplifies to:

$$C_{ab}(\nu) = \sum_{\vec{S}} S_a S_b p_{\vec{S}}(\nu) = C(\nu). \quad (4.31)$$

Moreover, the probability attached to any configuration $\vec{S}$ can be written as:

$$p_{\vec{S}}(\nu) = \frac{1}{2} \left(p_{\vec{S}}^-(\nu) + p_{\vec{S}}^+(\nu) \right), \quad (4.32)$$

where $p_{\vec{S}}^-(\nu)$ and $p_{\vec{S}}^+(\nu)$ are the probabilities of configuration $\vec{S}$, conditional on whether it is generated for group -1 or for group $+1$ respectively. In turn, these probabilities can be factorized in terms of feature-level probabilities of traits:

$$p_{\vec{S}}^-(\nu) = \prod_{a=1}^F p_{S_a}^-(\nu), \quad p_{\vec{S}}^+(\nu) = \prod_{a=1}^F p_{S_a}^+(\nu), \quad (4.33)$$

because once the group is chosen, each trait S_a (with possible values -1 and $+1$) is chosen independently at the level of the respective feature a . By inserting Eq. (4.33) in Eq. (4.32) and the result in Eq. (4.31), by carrying out appropriate algebraic manipulations, while making use of the fact that $\sum_{\vec{S}} = \prod_{a=1}^F (\sum_{S_a})$ and of the fact that $p_{S_a=-1}^{-/+}(\nu) + p_{S_a=+1}^{-/+}(\nu) = 1.0$, one obtains:

$$C(\nu) = \frac{1}{2} [p_{--}^-(\nu) - p_{-+}^-(\nu) - p_{+-}^-(\nu) + p_{++}^-(\nu)] + \frac{1}{2} [p_{--}^+(\nu) - p_{-+}^+(\nu) - p_{+-}^+(\nu) + p_{++}^+(\nu)], \quad (4.34)$$

where, for instance, $p_{-+}^-(\nu)$ is the probability that trait -1 is chosen for one of the features and that trait $+1$ is chosen for the other feature, conditional on the given configuration being generated for group -1 . Based on the model description in Sec. 4.3.2, one can see that:

$$p_{--}^-(\nu) = p_{++}^+(\nu) = (1 - 2\nu)^2, \quad (4.35)$$

$$p_{++}^-(\nu) = p_{--}^+(\nu) = (2\nu)^2, \quad (4.36)$$

$$p_{-+}^-(\nu) = p_{+-}^+(\nu) = (1 - 2\nu)(2\nu), \quad (4.37)$$

$$p_{+-}^-(\nu) = p_{-+}^+(\nu) = (2\nu)(1 - 2\nu). \quad (4.38)$$

By plugging these in Eq. (4.34), after simple algebraic manipulations, one obtains Eq. 4.10.

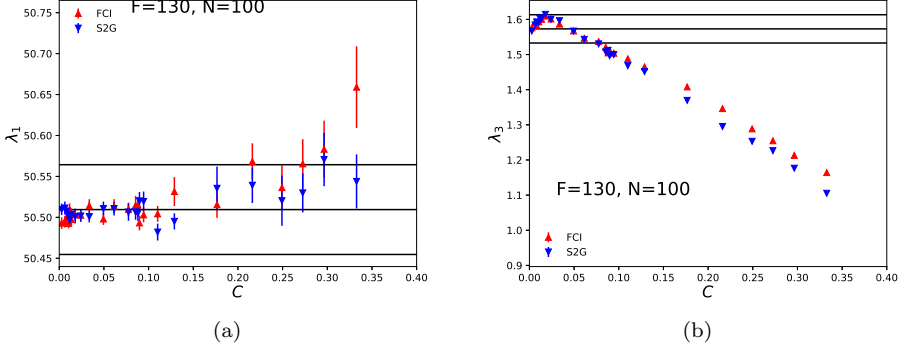


Figure 4.13: Behaviour of largest and third-largest eigenvalues λ_1 and λ_3 . The figure shows how λ_1 (a) and λ_3 (b) depend on the correlation level C for the fully-connected Ising (FCI, red, upward triangles) and for the symmetric two-groups (S2G, blue, downward triangles) models. For each C value, for each of the two models, an averaging is performed over 80 sets of cultural vectors independently sampled from the respective ensemble – the vertical bar associated to each point shows the interval spanned by one standard mean error on each side of the mean. The black, horizontal lines show, for comparison, the mean λ_1 (a) and mean λ_3 (b) based on uniform randomness, along with the width of the λ_1 and the λ_3 distributions – one standard deviation on each side – where the calculations are based on 60 sets of cultural vectors generated via uniform randomness – these lines do not imply that, for uniform randomness, the correlation C (which actually vanishes by construction) can be arbitrarily large.

4.C The structure of the FCI and S2G models

This section shows that the structure implicit in cultural states generated with either the FCI or the S2G model is captured by only one eigenpair of the similarity matrix, so that there is at most one structural mode. Specifically, as the correlation level is increased for the FCI and the S2G models, there is only one eigenvalue – the subdominant eigenvalue λ_2 – that becomes separated from the random bulk, while becoming significantly larger than the upper boundary of the bulk that is expected based on uniform randomness. The behavior of λ_2 has already been presented in Fig. 4.8. The results shown here, via Fig. 4.13, are complementary to those shown in Fig. 4.8, which uses the same format, while focusing on the behavior of λ_1 in Fig. 4.13(a) and on the behavior of λ_3 in Fig. 4.13(b). Note that λ_1 , associated to the global mode, remains statistically compatible with the null model as the level of correlation is increased, for both FCI and S2G. On the other hand, λ_3 decreases, while becoming, for large enough C , significantly smaller than the upper boundary of the bulk predicted by uniform randomness. All this shows

that the structure FCI and S2G is mostly captured by the eigenpair of λ_2 , which becomes increasingly stronger as the correlation level increases. This appears to be a consequence of the fact that each model is controlled by one parameter, while all the non-uniformity of the associate probability distribution is captured by one dimension, namely the F_+ axis of Fig. 4.6.

Bibliography

- [1] John Urry. The complexity turn. *Theory, Culture & Society*, 22(5):1–14, 2005.
- [2] David Lazer, Alex Pentland, Lada Adamic, Sinan Aral, Albert-László Barabási, Devon Brewer, Nicholas Christakis, Noshir Contractor, James Fowler, Myron Gutmann, Tony Jebara, Gary King, Michael Macy, Deb Roy, and Marshall Van Alstyne. Computational social science. *Science*, 323(5915):721–723, 2009.
- [3] Charles Kadushin. *Understanding Social Networks: Theories, Concepts and Findings*. Oxford University Press, 2012.
- [4] Claudio Castellano, Santo Fortunato, and Vittorio Loreto. Statistical physics of social dynamics. *Rev. Mod. Phys.*, 81:591–646, May 2009.
- [5] Luca Valori, Francesco Picciolo, Agnes Allansdottir, and Diego Garlaschelli. Reconciling long-term cultural diversity and short-term collective social behavior. *Proc. Natl. Acad. Sci.*, 109(4):1068–1073, 2012.
- [6] Alex Stivala, Garry Robins, Yoshihisa Kashima, and Michael Kirley. Ultrametric distribution of culture vectors in an extended Axelrod model of cultural dissemination. *Sci. Rep.*, 4(4870), 2014.
- [7] Alexandru-Ionuț Băbeanu, Leandros Talman, and Diego Garlaschelli. Signs of universality in the structure of culture. *The European Physical Journal B*, 90(12):237, Dec 2017.
- [8] Alexandru-Ionuț Băbeanu and Diego Garlaschelli. Evidence for mixed rationalities in preference formation. *Complexity*, 2018, 2018. Article ID 3615476.
- [9] Alexandru-Ionuț Băbeanu, Jorinde van de Vis, and Diego Garlaschelli. Ultrametricity increases the predictability of cultural dynamics. *arXiv:1712.05959v1*, 2017.
- [10] Robert Axelrod. The dissemination of culture. *Journal of Conflict Resolution*, 41(2):203–226, 1997.
- [11] Madan Lal Mehta. *Random Matrices*. Elsevier, 2012.

- [12] Alan Edelman and N. Raj Rao. Random matrix theory. *Acta Numerica*, 14:233–297, 2005.
- [13] Marc Potters, Jean-Philippe Bouchaud, and Laurent Laloux. Financial applications of random matrix theory: old laces and new pieces. *Acta Phys. Pol. B*, 36(9):2767–2784, 2005.
- [14] Mel MacMahon and Diego Garlaschelli. Community detection for correlation matrices. *Phys. Rev. X*, 5:021006, Apr 2015.
- [15] Assaf Almog, Ori Roethler, Renate Buijink, Stephan Michel, Johanna H Meijer, Jos H. T. Rohling, and Diego Garlaschelli. Uncovering functional brain signature via random matrix theory. *arXiv:1708.07046v2*, 2017.
- [16] V. A. Marchenko and L. A. Pastur. Distribution of eigenvalues for some sets of random matrices. *Math. USSR-Sb.*, 1:457–483, 1967.
- [17] Aashay Patil and M. S. Santhanam. Random matrix approach to categorical data analysis. *Phys. Rev. E*, 92:032130, Sep 2015.
- [18] E. T. Jaynes. Information theory and statistical mechanics. *Phys. Rev.*, 106:620–630, May 1957.
- [19] Louis Colonna-Romano, Harvey Gould, and W. Klein. Anomalous mean-field behavior of the fully connected ising model. *Phys. Rev. E*, 90:042111, Oct 2014.
- [20] Michael Thompson, Richard J. Ellis, and Aaron Wildavsky. *Cultural Theory*. Westview Press, 1990.
- [21] Karlheinz Reif and Anna Melich. Euro-barometer 38.1: Consumer protection and perceptions of science and technology, november 1992. <http://www.icpsr.umich.edu/icpsrweb/ICPSR/studies/06045>, 1995.
- [22] Tom W. Smith, Peter Marsden, Michael Hout, and Jibum Kim. General social surveys, 1993 ed. <http://gss.norc.org/get-the-data/spss>, 1972–2012.
- [23] Ken Goldberg, Theresa Roeder, Dhruv Gupta, and Chris Perkins. Eigen-taste: A constant time collaborative filtering algorithm. *Information Retrieval*, 4(2):133–151, 2001.
- [24] K R W Jones. Entropy of random quantum states. *Journal of Physics A: Mathematical and General*, 23(23):L1247, 1990.
- [25] George H Dunteman. *Principal components analysis*. Sage Publications, 1989.
- [26] Leonard Kaufman and Peter J. Rousseeuw. *Finding groups in data: An Introduction to Cluster Analysis*. John Wiley & Sons, 1990.

Samenvatting

Samenlevingen vertonen meerdere structurele en dynamische eigenschappen die interessant zijn vanuit het perspectief van complexe systemen. Zulke eigenschappen kunnen op twee manieren worden benaderd. Bij de sociale benadering gaat het om de sociale netwerken, waarin verbindingspatronen tussen individuen bekeken worden. Bij de culturele benadering gaat het om de verscheidenheid aan meningen, voorkeuren en andere culturele attributen onder de individuen. Binnen de sociale benadering is er zowel datagedreven (empirisch), als modelgedreven (theoretisch) onderzoek uitgevoerd, waarbij de nadruk ligt op zowel de structuur als de dynamiek van sociale netwerken. Dit in tegenstelling tot de culturele benadering, waarbij voornamelijk theoretisch onderzoek is uitgevoerd, met weinig tot geen empirische input. Dan gaat het vooral om de manier waarop de culturele karakteristieken veranderen door de onderlinge beïnvloeding van de individuen. Dit proefschrift doet een stap voorwaarts om deze balans te verbeteren, door de focus te leggen op de structurele eigenschappen van de cultuur, zoals gevangen door statische, multidimensionale, empirische data uit grootschalige sociale enquêtes.

Als eerste stap worden deze data omgezet naar een symbolische keten van culturele attributen, bekend als “culturele vectoren”, die geassocieerd worden met verschillende individuen. Hier geldt dat verschillende posities in elke keten overeenkomen met verschillende vragen van de enquête. Verschillende empirische bronnen worden gebruikt om meerdere sets van culturele vectoren op te bouwen, die bekend staan als “culturele toestanden”. Deze worden vervolgens geanalyseerd met een eerder ontwikkelde techniek die twee grootheden combineert die beiden onafhankelijk zijn van de preciese enquêtevragen: de neiging tot culturele diversiteit op lange termijn en de neiging tot sociale coördinatie op korte termijn. Beiden zijn gebaseerd op de theorie van de dynamica van sociale invloeden. Deze techniek bevat ook een vergelijking tussen de empirische data en geschikte, gerandomiseerde, tegenhangers. Deze analyse legt verschillen bloot tussen de empirische culturele toestanden en de gerandomiseerde tegenhangers, alsmede opvallende overeenkomsten tussen verschillende datasets. Dit suggereert dat er niet-triviale, universele eigenschappen ten grondslag liggen aan de structuur van culturele eigenschappen.

Als tweede stap is het mechanisme achter de robuuste empirische eigenschappen onderzocht. Dit leidt tot het voorstel van een statisch probabilistisch model,

dat in staat is tot het genereren van culturele toestanden die deze eigenschappen reproduceren. Het model neemt aan dat iedere individuele reeks van attributen gedeeltelijk bepaald wordt door een of meerdere, vermoedelijk universele, “rationaliteiten”, wier bestaan informeel wordt aangenomen door meerdere sociaal-wetenschappelijke theorieën. Verder neemt het model aan dat, naast een dominante rationaliteit, ieder individu ook affiniteit heeft met de andere rationaliteiten. Er wordt aangetoond dat beide aannames nodig zijn voor het reproduceren van empirische regelmatigheden. Dit impliceert dat de generieke structuur van cultuur in overeenstemming is met het bestaan van meerdere, gemixte rationaliteiten, wat indirect bewijs levert voor sociaal-wetenschappelijke theorieën gebaseerd op dit idee.

Als derde stap bekijkt dit proefschrift de beperkingen die de empirische structuur legt op de culturele dynamica op lange termijn, die wordt aangedreven door sociale invloeden. Precieser: er wordt berekend in hoeverre de samenstelling van de groepen in de eindtoestand, geproduceerd door een simpel model van culturele dynamica, voorspeld kan worden door de culturele vectoren die de begintoestand specificeren, zonder de dynamica expliciet uit te voeren. Het is aangetoond dat de voorspelbaarheid, die rigoureus is gedefinieerd in een informatie-theoretische zin, significant hoger is voor empirische culturele toestanden dan voor gerandomiseerde tegenhangers vanwege de hiërarchische ultrametrick-achtige organisatie van de eerste, die de culturele convergentie beperkt tot de lagere niveaus van de hiërarchie. Bovendien gaat hogere voorspelbaarheid hand in hand met hogere compatibiliteit van de korte-termijn sociale coördinatie en lange-termijn culturele diversiteit, die een essentieel aspect is van de bovenvermelde empirische robuustheid. Verder is een nulmodel geïntroduceerd om culturele begintoestanden te genereren die de ultrametrick-vorm van de data behouden. Gebruikmakend van dit ultrametrick-model wordt de voorspelbaarheid sterk verbeterd in vergelijking met de gerandomiseerde tegenhangers. Dit bevestigt dat de hiërarchische organisatie van werkelijke cultuur zeer belangrijk is voor het voorspellen van het resultaat van de dynamica van sociale invloeden.

Als vierde en laatste stap wordt de structuur die inherent is aan empirische culturele toestanden nader onderzocht, met behulp van concepten uit de random matrix theorie, toegepast op similariteitsmatrices tussen culturele vectoren. Voor het genereren van random matrices die geschikt zijn als structuurloze referentie, stellen we een nulmodel voor dat de empirische frequentie van elk mogelijk cultureel attribuut gemiddeld oplegt. Met betrekking tot dit nulmodel vertonen de empirische similariteitsmatrices afwijkende eigenwaarden, die mogelijk een teken zijn van culturele groepen of clusters die mogelijk niet op andere manieren herkenbaar zijn. Ze kunnen echter ook artefacten zijn van willekeurige, datasetafhankelijke correlaties tussen culturele variabelen. Deze mogelijkheid wordt expliciet geïllustreerd met behulp van twee eenvoudige modellen. In het eerste wordt het “groepen-scenario” geïmplementeerd en in het tweede het “correlatie-scenario”. Ook wordt in deze setting aangetoond dat de twee scenario’s kunnen worden onderscheiden door de uniformiteit van de waarden van de componenten van de

eigenvector te berekenen die bij een afwijkende eigenwaarde hoort. Tegelijk wordt gecontroleerd of deze uniformiteit statistisch overeenkomt met het nulmodel. Voor empirische data wordt aangetoond dat de eigenvectoruniformiteiten van alle afwijkende eigenwaarden verenigbaar zijn met het nulmodel. Dit suggereert dat de ogenschijnlijke groepsstructuur niet echt is. Afwijkende eigenvectoruniformiteiten zouden echter afwezig kunnen zijn voor culturele groepen die worden veroorzaakt door het mengen van rationaliteiten (de plausibele structurele hypothese hierboven genoemd). Verder onderzoek is daarom nodig om een beslissende uitspraak over de aanwezigheid of afwezigheid van groepsstructuur in cultuur te doen.

Summary

Human societies exhibit a multitude of structural and dynamical properties that are interesting from a complex systems perspective. Such properties can be identified at two levels of analysis. On one hand, one is confronted with social networks, capturing patterns of connectivity and interactions between social agents – the “social” level of analysis. On the other hand, one is confronted with distributions of opinions, preferences and other cultural traits among the agents – the “cultural” level of analysis. The social level has seen a healthy mix of empirical, data-driven research and of theoretical, model-driven research, focusing on both the structure and the dynamics of social networks. By contrast, the cultural level has mostly seen theoretical, model-driven research, with little or no empirical input, with a strong emphasis on the dynamics of cultural traits, driven by social influence interactions between agents. This thesis can be seen as a step towards compensating for this imbalance, as it focuses on structural properties of culture, captured by static, multidimensional empirical data from large-scale social surveys.

As a first step, this data is converted into symbolic sequences of cultural traits, known as “cultural vectors,” associated to different individuals, where different positions in each sequence correspond to different survey questions. Different empirical sources are used for constructing multiple sets of cultural vectors, where one such set is also called a “cultural state.” These are analyzed with a previously developed technique, which combines two quantities whose definitions are independent of the set of survey questions: a measure of propensity to long-term cultural diversity and a measure of propensity to short-term social coordination, both of which are based on theoretical notions of social influence dynamics. The technique also incorporates a comparison between empirical data and appropriate randomized counterparts. The analysis shows clear deviations of empirical cultural states from randomized counterparts, as well as remarkable similarities across different datasets, suggesting that there are non-trivial, universal properties underlying the structure of culture.

As a second step, the mechanism behind the robust empirical properties is investigated. This leads to proposing a static, probabilistic model capable of generating cultural states that reproduce these properties. The model assumes that every individual’s sequence of traits is partly dictated by one of several supposedly

universal “rationalities,” whose existence is informally postulated by several social science theories. In addition, the model assumes that, apart from a dominant rationality, each individual also has some affinity with the other rationalities. It is shown that both assumptions are required for reproducing the empirical regularities. This implies that the generic structure of culture is compatible with the existence of several, mixing rationalities, providing indirect evidence for social science theories that are based on this idea.

As a third step, this thesis examines the constraints that empirical structure places on long-term cultural dynamics driven by social influence. More precisely, it evaluates the extent to which the contents of the final state groups (the subsets of agents whose cultural vectors are identical in the final state), produced by a simple model of cultural dynamics, can be predicted based on the cultural vectors that specify the initial cultural state, without explicitly running the dynamics. This predictability, which is rigorously defined in an information-theoretic sense, is shown to be significantly higher for empirical cultural states than for randomized counterparts, due to the hierarchical ultrametric-like organization of the former, which confines cultural convergence within the lower levels of the hierarchy. Moreover, higher predictability goes along with higher compatibility of short-term social coordination and long-term cultural diversity, which is an essential aspect of the empirical robustness mentioned above. In addition, a null model is introduced for generating initial cultural states that retain the ultrametric representation of real data. Using this ultrametric model, predictability is highly enhanced with respect to the randomized cases. This confirms the importance of the hierarchical organization of real culture for forecasting the outcome of social influence dynamics.

As a fourth and final step, the structure inherent in empirical cultural states is further investigated, using concepts from random matrix theory, applied to matrices of similarity between cultural vectors. For generating random matrices that are appropriate as a structureless reference, we propose a null model that enforces, on average, the empirical occurrence frequency of each possible trait. With respect to this null model, the empirical similarity matrices show deviating eigenvalues, which may be signatures of cultural groups or clusters that might not be recognizable by other means. However, they can conceivably also be artifacts of arbitrary, dataset-dependent correlations between cultural variables. This possibility is explicitly illustrated, with the help of two toy models, which implement the “groups scenario” and the “correlations scenario” respectively, in the simplest conceivable setting. It is also shown that, at least in this setting, the two scenarios can be distinguished by evaluating the uniformity of the entries of the eigenvector associated to a deviating eigenvalue, while checking if this uniformity is statistically compatible with the null model. For empirical data, the eigenvector uniformities of all deviating eigenvalues are shown to be compatible with the null model, suggesting that the apparent group structure is not genuine. However, deviating eigenvector uniformities might not be present for cultural groups induced by mixing rationalities (the plausible structural hypothesis mentioned above), so

further research is required for a decisive statement about the presence or absence of group structure in culture.

Rezumat

Societățile umane prezintă o serie de proprietăți dinamice și structurale care sunt interesante din perspectiva sistemelor complexe. Astfel de proprietăți pot fi identificate pe două niveluri de analiză. Pe de o parte, se poate vorbi despre rețele sociale, care încorporează regularități de conectivitate și interacțiuni între agenți sociali – nivelul “social” de analiză. Pe de altă parte, se poate vorbi despre distribuții de opinii, preferințe și alte trăsături culturale în populația de agenți – nivelul “cultural” de analiză. În cadrul nivelului social, există atât cercetare empirică, pe bază de date, cât și cercetare teoretică, pe bază de modele matematice, cercetare care se axează atât pe structura cât și pe dinamica rețelelor sociale. Prin antiteză, în cadrul nivelului cultural, cercetarea este predominant teoretică, bazată pe modele matematice, cu foarte puține contribuții empirice, cercetare axată preponderent pe dinamica trăsăturilor culturale, datorată interacțiunilor de influență socială între agenți. Această teză poate fi văzută ca un pas înainte către corectarea acestui dezechilibru, deoarece se concentrează pe proprietățile structurale ale culturii, accesibile prin date empirice multidimensionale obținute prin sondaje de opinie la scară largă.

Ca un prim pas, aceste date sunt convertite în secvențe simbolice de trăsături culturale, denumite “vectori culturali”, asociați persoanelor participante la respectivul sondaj de opinie, unde o anumită poziție într-o secvență corespunde unei anumite întrebări din chestionar. Diferite surse empirice sunt folosite pentru a construi mai multe astfel de seturi de vectori culturali, unde un astfel de set este denumit și o “stare culturală”. Acestea sunt analizate cu ajutorul unei tehnici dezvoltate anterior, care combină două cantități ale căror definiții sunt independente de setul de întrebări din chestionar: prima cantitate masoară predilecția către diversitate culturală pe termen lung, în timp ce a doua cantitate masoară predilecția către coordonare socială pe termen scurt, ambele bazându-se pe noțiuni teoretice ale dinamicii sub influență socială. Această tehnică încorporează și o comparație între datele empirice și date randomizate asociate. Analiza arată o deviație clară a stărilor culturale empirice față de stările randomizate asociate, precum și similități remarcabile între diferite seturi de date, sugerând că structura culturală a societăților reale este caracterizată de proprietăți netriviiale și universale.

Ca un al doilea pas, se investighează mecanismul din spatele proprietăților empirice robuste. Ca urmare, un model probabilistic este propus, capabil să genereze

stări culturale care reproduc aceste proprietăți. Modelul presupune că secvența de trăsături a fiecărei persoane este dictată parțial de către o “raționalitate”, dintr-un set restrâns de raționalități, presupus universale, a căror existență este postulată informal de o serie de teorii aparținând științelor sociale. În plus, modelul presupune că, pe lângă o raționalitate dominantă, fiecare persoană are o oarecare afinitate și cu celelalte raționalități. Se arată că ambele presupuneri sunt necesare pentru a reproduce regularitățile empirice. Aceasta înseamnă că structura generică a culturii este compatibilă cu existența mai multor raționalități, combinate și integrate în mod diferit la nivelul fiecărui individ, ceea ce furnizează dovezi indirecte pentru teorii sociologice care se bazează pe această idee.

Ca un al treilea pas, această teză examinează constrângerile pe care structura empirică le impune asupra dinamicii culturale pe termen lung, sub acțiunea influenței sociale. Mai exact, se evaluează gradul în care conținutul grupurilor din starea finală (seturi de agenți ai căror vectori culturali sunt identici în starea finală), produsă de un model simplu de dinamică culturală, poate fi prezis pe baza vectorilor culturali care specifică starea inițială, fără a rula dinamica în mod explicit. Se arată că această predictibilitate, definită riguros prin noțiuni ale teoriei informației, este semnificativ mai mare pentru stările culturale empirice decât pentru omoloagele lor randomizate, datorită organizării ierarhic-ultrametrice a celor dintâi, organizare ce menține convergența culturală preponderent înăuntrul nivelurilor inferioare ale ierarhiei. Mai mult decât atât, un nivel ridicat de predictibilitate se asociază unui grad ridicat de compatibilitate între coordonarea socială pe termen scurt și diversitatea culturală pe termen lung, un aspect definitoriu al structurii empirice menționate anterior. În plus, un nou model probabilistic este prezentat, care generează stări culturale inițiale pe baza reprezentării ultrametrice a datelor reale. Stările culturale generate pe baza acestui model ultrametric prezintă o predictibilitate semnificativ mărită față de omoloagele randomizate. Astfel, se confirmă importanța organizării ierarhice a culturii pentru prognoza dinamicii de influență socială în sisteme reale.

Ca un ultim pas, structura inerentă stărilor culturale empirice este investigată mai detaliat, cu ajutorul unor noțiuni ale teoriei matricilor aleatorii, aplicate matricilor de similaritate dintre vectorii culturali. Pentru a genera matrici aleatoare adecvate ca un scenariu de referință lipsit de structură, propunem un model nul (model probabilistic bazat pe o ipoteză nulă) care reproduce, în medie, frecvența empirică de apariție a oricărei trăsături posibile. Prin comparație cu acest model nul, matricile empirice de similaritate prezintă valori proprii ce deviază semnificativ, ceea ce poate semnala prezența unor grupuri ce ar putea fi nedetectabile prin alte metode. Totuși, este posibil ca prezența acestor valori proprii deviate să fie o simplă consecință a unor corelații între variabilele culturale, corelații arbitrar, specifice fiecărui set de date. Această posibilitate este ilustrată în mod explicit, cu ajutorul a două modele simpliste, care exemplifică “scenariul grupurilor”, respectiv “scenariul corelațiilor”, în cea mai simplă manieră posibilă. De asemenea, se arată că, cel puțin în aceste condiții artificiale, controlate, cele două scenarii pot fi distinse cu ajutorul uniformității elementelor vectorului propriu asociat unei valori

propriu deviate, verificând dacă această uniformitate este compatibilă statistic cu modelul nul. Se arată că, pentru datele empirice, vectorii proprii asociați tuturor valorilor proprii deviate sunt compatibili cu modelul nul din punctul de vedere al uniformității, sugerând că nu avem de-a face, de fapt, cu grupuri culturale autentice. Cu toate acestea, grupurile culturale induse de raționalități combinate (ipoteza structurală plausibilă menționată anterior) ar putea, la rândul lor, să nu prezinte uniformități deviate ale vectorilor proprii. Drept urmare, mai multă cercetare este necesară pentru a stabili, în mod decisiv, prezența sau absența unor grupuri autentice în structura culturii.

List of publications

- [1] *Stability of superconducting strings coupled to cosmic strings*, A. Babeanu and B. Hartmann, *Physical Review D* **85** 2 023518 (2012).
- [2] *Observation of the $\Lambda_b^0 \rightarrow J/\psi p \pi^-$ decay*, The LHCb collaboration, *Journal of High Energy Physics* **2014** 7 103 (2014).
- [3] *Signs of universality in the structure of culture*, A. I. Băbeanu, L. Talman and D. Garlaschelli, *European Physics Journal B* **90** 12 237 (2017).
- [4] *Evidence for mixed rationalities in preference formation*, A. I. Băbeanu and D. Garlaschelli, *Complexity* **2018** 3615476 (2018).
- [5] *Ultrametricity increases the predictability of cultural dynamics*, A. I. Băbeanu, J. van de Vis and D. Garlaschelli, arXiv:1712.05959 (2017).
- [6] *A random matrix perspective of cultural structure*, A. I. Băbeanu, arXiv:1803.04324 (2018).

Curriculum vitæ

I was born in Constanța, Romania, on the 28th of September 1989. In Constanța I also completed my pre-university education, upon graduating, in June 2008, from the “Mircea cel Bătrân” National College (high school), with a specialization in computer science and mathematics. I continued my education at Jacobs University Bremen (Germany), where I was awarded a Bachelor’s of Science in physics, in June 2011. I moved to the Netherlands in September 2011, after being awarded a Huygens scholarship, which fully covered the tuition and living expenses for my Master’s education in particle physics at Utrecht University. In September 2013 I started my Ph.D. education at Leiden University, funded by a Leiden/Huygens fellowship in physics, in the group of Dr. Diego Garlaschelli. The research conducted there is described in this thesis. Parts of this research have been presented at several scientific conferences and workshops, in Finland, Italy, Germany and The Netherlands. In parallel with my research, I attended several specialized schools in The Netherlands and France, while also serving as a teaching assistant for the “Molecular Physics for Life Science and Technology” course and as an organizer of the “Leiden Complex Networks Network” (LCN2) initiative. Starting with September 2018, I am working as a software/web developer at DongIT in Leiden.

Acknowledgments

First of all, I am very grateful to my supervisor, Diego Garlaschelli, for his guidance, support and patience during years of challenging research. In addition to the insightful scientific collaborations and discussions, interacting with him taught me how to maintain a balance between rigor, creativity and accessibility in my work. The lessons I have learned from him will have a profound effect on my career, regardless of my future choices.

My research greatly benefited from the involvement of Arjen Aerts, Leandros Talman and Jorinde van de Vis, whose projects I had the pleasure to supervise. It also benefited from several rounds of thoughtful comments of Maroussia Favre, as well as from exciting interactions with Marco Verweij and Michael Thompson, about fundamental aspects of social science. In addition, I acknowledge constructive discussion with Assaf Almog, Ulf Dieckmann, Andreas Flache, Santo Fortunato, Elena Garuccio, Valerio Gemmetto, Gerard 't Hooft, Michael Mäs, Petter Säterskog, Frank Takes and Vincent Traag. I am also grateful to Peter Denteneer, Janusz Meylahn, Jorgos Papadomanolakis, Corrie Wortel and Renate van Zanten for helping with the Dutch translation of the Summary.

In addition to people already mentioned above, I would like to thank my other colleagues and friends from the Econophysics and Network Theory group: Alex, Andrea, Elena, Federica, Ferry, Maria, Marc, Nedim, Qi, Ruben, Stella, Tiziano and Vasyl, as well as from the encompassing Physics Institute: Alexey, Anastasia, Andrii, Artem, Cameron, Eric, Ireth, Koenraad, Kyrylo, Marco, Marco, Maksym, Martina, Nicandro, Oksana, Paul, Richard, Theo, Vincenzo, Yaroslav and Yvette, for their warmth, company and support. I am also thankful to the administrative staff of our Physics Institute, in particular to Barry, Fran, Marianne, Manon and Monique, for providing fast solutions to a variety of problems.

Last but not least, I am grateful to my other friends in Leiden and beyond, as well as to my family, for their continuous support over the years.