



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

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.

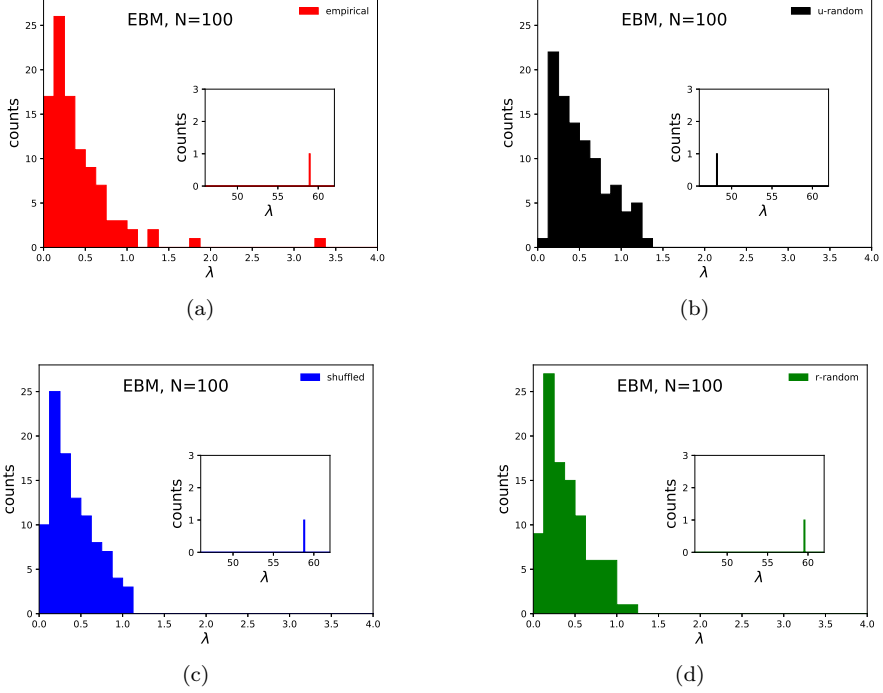


Figure 4.1: Eigenvalue spectra of cultural similarity matrices. The first panel correspond to an empirical cultural state (a) with $N = 100$ vectors constructed from Eurobarometer (EBM) data, while the other three correspond to associated cultural states generated with the uniform random (b), the shuffled (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

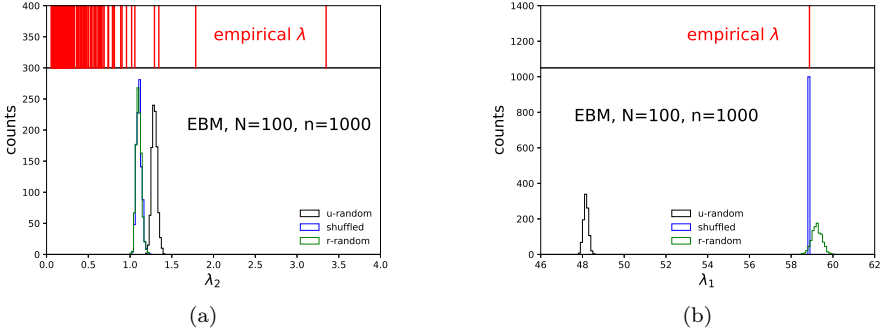


Figure 4.2: Leading and subleading eigenvalue distributions for random matrices. The figure shows the subleading eigenvalue λ_2 distribution (a) and the leading eigenvalue λ_1 distribution (b), for the three null models (legends), implementing uniform randomness (black), shuffling (blue) and restricted 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

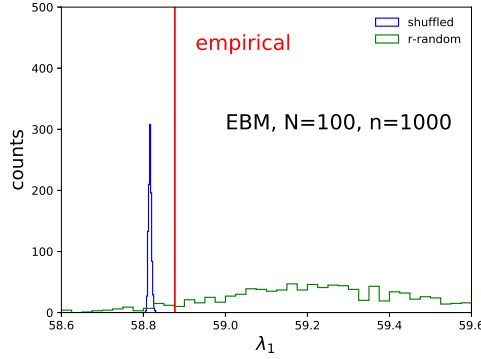


Figure 4.3: Detailed comparison in terms of the leading eigenvalue. The figure shows in detail the distributions of the leading eigenvalue λ_1 for the shuffled (blue) and restricted random (green) null models, in comparison to the empirical value (vertical red line). For each distribution, $n = 1000$ 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.

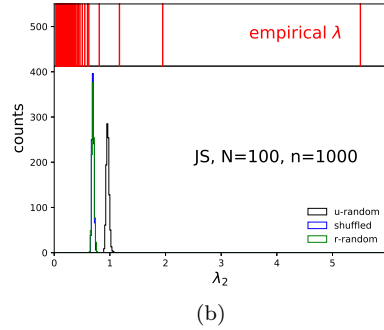
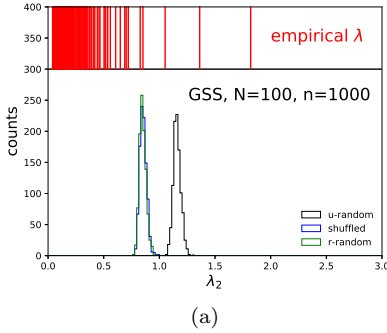


Figure 4.4: Empirical structure in other datasets. The figure shows the subleading eigenvalue λ_2 distribution for the three null models (legends), implementing uniform randomness (black), shuffling (blue) and restricted randomness (green), in comparison to the empirical eigenvalues, whose positions 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 $C(0, F) = 0.0$ 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.

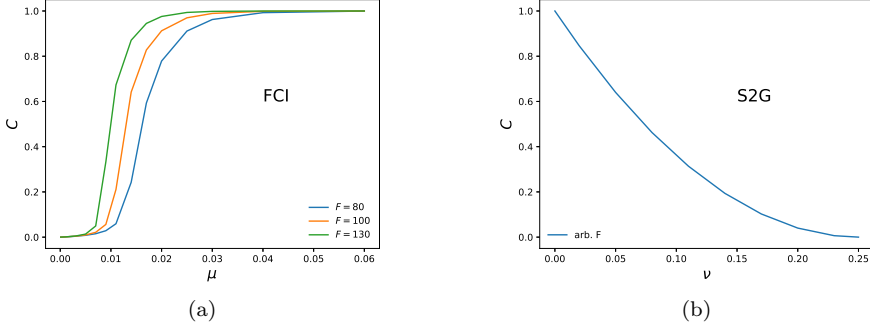


Figure 4.5: Correlation behaviour. The figure shows the dependence of the pairwise feature correlation C , first (a) on the feature-feature coupling strength parameter μ controlling the fully-connected Ising model (FCI), second (b) on the group strength parameter ν controlling the symmetric two-groups 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

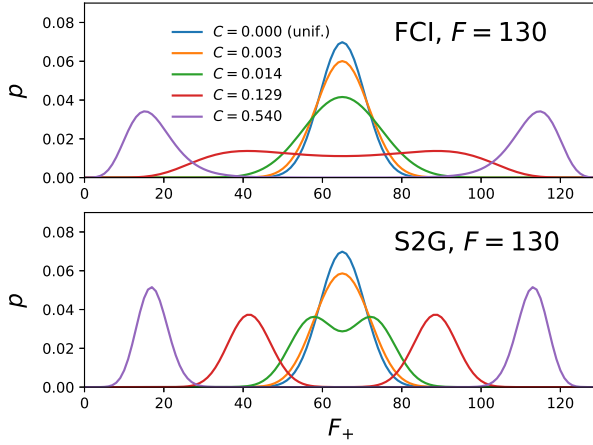


Figure 4.6: Shape of probability distribution. The figure shows the probability associated to configurations with F_+ traits +1 for $F = 130$ features, for different values of the feature-feature correlation level C (legend), for the fully-connected Ising model (FCI, top) and for the symmetric two-groups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

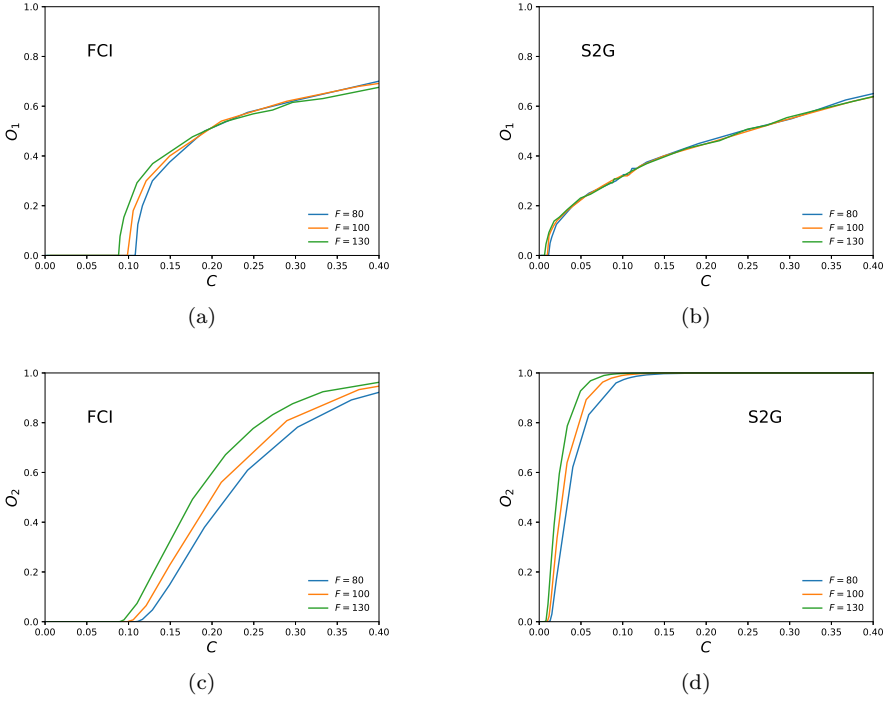


Figure 4.7: Symmetry breaking phase transitions. The figure shows the behavior of the normalized probability peak departure O_1 (top) and of the normalized peak height O_2 (bottom), as functions of the model parameters, for the fully-connected Ising (FCI, top) and the symmetric two-groups (S2G, bottom) 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

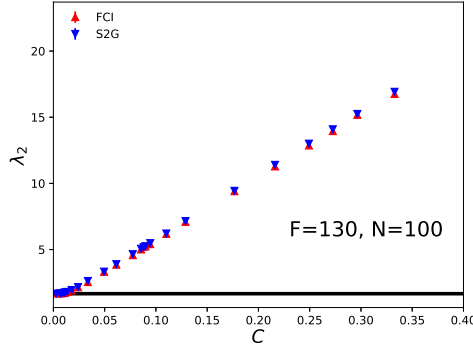


Figure 4.8: Behaviour of subleading eigenvalue (λ_2). The figure shows how λ_2 depends on the correlation level C for the fully-connected Ising (FCI, red, upward triangles) and for the symmetric two-groups (S2G, blue, downward triangles) models. For each C value, for each of the two models, an averaging 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-

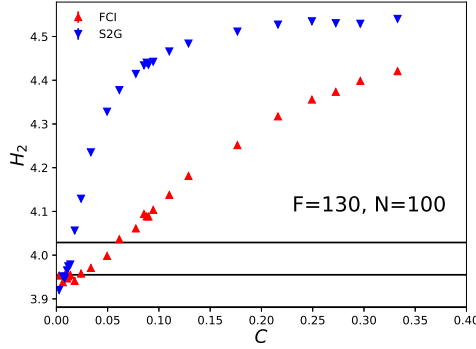


Figure 4.9: Behaviour of uniformity H_2 associated to subleading eigenvalue. The figure shows how H_2 depends on the correlation level C for the fully-connected Ising (FCI, red) and for the symmetric two-groups (S2G, blue) models. For each C value, for each of the two models, an averaging is performed 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

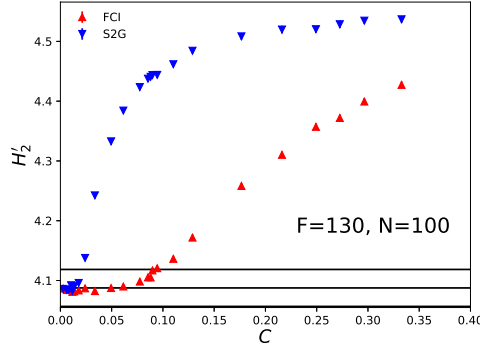


Figure 4.10: Behavior of subleading uniformity H'_2 . The figure shows how H'_2 depends on the correlation level C for the fully-connected Ising (FCI, red) and for the symmetric two-groups (S2G, blue) models. For each C value, for each of the two models, an averaging is performed over 80 sets of cultural 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

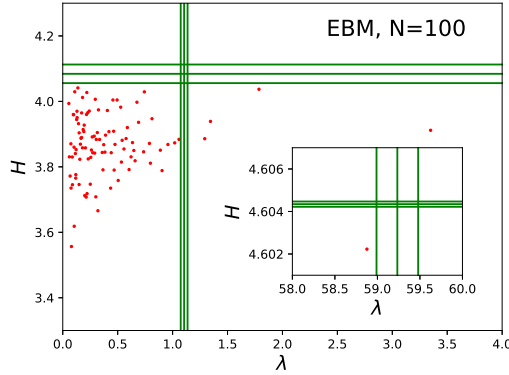


Figure 4.11: Eigenvalues and eigenvector entropies for empirical data. Every point corresponds to an empirical eigenpair, with the eigenvalue λ shown along the horizontal axis and the eigenvector entropy H shown along the vertical axis. The inset focuses on the leading eigenvalue, which also corresponds 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.

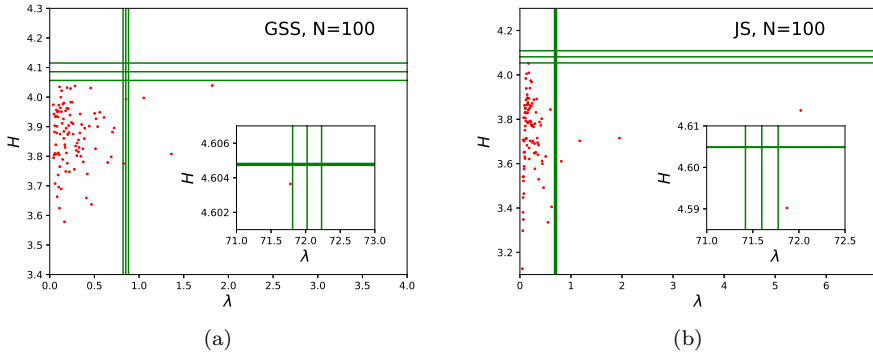


Figure 4.12: Eigenvalues and eigenvector entropies in other empirical datasets. Fig. (a) is based on the General Social Survey (GSS) data with $N = 100$ cultural vectors used for Fig. 4.12(a), while Fig. (b) is based on the Jester (JS) data with $N = 100$ cultural vectors used in Figs. 4.12(b). In each 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.

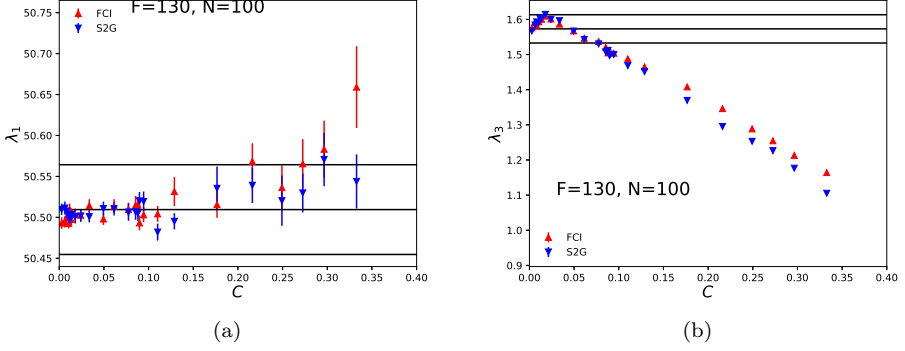


Figure 4.13: Behaviour of largest and third-largest eigenvalues λ_1 and λ_3 . The figure shows how λ_1 (a) and λ_3 (b) depend on the correlation level C for the fully-connected Ising (FCI, red, upward triangles) and for the symmetric two-groups (S2G, blue, downward triangles) models. For each C value, 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.

