



Universiteit
Leiden
The Netherlands

Reproducing the first and second moments of empirical degree distributions

Mattia, M.; Giuffrida, F.; Garlaschelli, D.; Squartini, T.

Citation

Mattia, M., Giuffrida, F., Garlaschelli, D., & Squartini, T. (2026). Reproducing the first and second moments of empirical degree distributions. *Physical Review Research*, 8(1).
doi:10.1103/3vtj-5nlt

Version: Publisher's Version

License: [Creative Commons CC BY 4.0 license](#)

Downloaded from: <https://hdl.handle.net/1887/4297271>

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

Reproducing the first and second moments of empirical degree distributions

Mattia Marzi ^{1,*}, Francesca Giuffrida ^{2,1,3}, Diego Garlaschelli ^{1,3,4} and Tiziano Squartini ^{1,5,4}

¹*IMT School for Advanced Studies, P.zza San Francesco 19, 55100 Lucca, Italy*

²*Dipartimento di Chimica e Fisica, Università di Palermo, V.le delle Scienze, Edificio 18, 90128 Palermo, Italy*

³*Lorentz Institute for Theoretical Physics, University of Leiden, Niels Bohrweg 2, 2333 CA Leiden, The Netherlands*

⁴*INdAM-GNAMPA Istituto Nazionale di Alta Matematica ‘Francesco Severi’, P.le Aldo Moro 5, 00185 Rome, Italy*

⁵*Scuola Normale Superiore, P.zza dei Cavalieri 7, 56126 Pisa, Italy*



(Received 31 July 2025; accepted 20 January 2026; published 9 February 2026)

The study of probabilistic models for the analysis of complex networks represents a flourishing research field. Among the former, exponential random graphs (ERGs) have gained increasing attention over the years. So far, only linear ERGs have been extensively employed to gain insight into the structural organization of real-world complex networks. None, however, is capable of accounting for the variance of the empirical degree distribution. To this aim, nonlinear ERGs must be considered. After showing that the usual mean-field approximation forces the degree-corrected version of the two-star model to degenerate, we define a fitness-induced variant of it. Such a “softened” model is capable of reproducing the sample variance, while retaining the explanatory power of its linear counterpart, within a purely canonical framework.

DOI: [10.1103/3vtj-5nlt](https://doi.org/10.1103/3vtj-5nlt)

I. INTRODUCTION

Network theory is employed to address problems of scientific and societal relevance, from the prediction of epidemic spreading to the identification of early-warning signals of upcoming financial crises [1–8]. As any dynamical process is strongly affected by the topology of the underlying network, one needs to individuate which higher-order properties can be traced back to lower-order ones and which, instead, are due to additional factors: This goal can be achieved by constructing ensembles of graphs whose defining properties are the same as in the real world but the topology is random under any other respect [9–15].

A class of models that has gained increasing attention over the years is that of exponential random graphs (ERGs) [10–18]. ERGs belong to the category of canonical approaches, being induced by constraints that are *soft*, i.e., constraints that can be violated by individual configurations even if their ensemble average matches the enforced value exactly. A key advantage of canonical approaches is that the expected value of topological properties can be often expressed analytically in terms of the constraints, thereby avoiding the computational cost of generating randomized networks [19]. Microcanonical approaches, instead, artificially generate many randomized variants of the observed network, enforcing constraints that are *hard*, i.e., constraints

that are met exactly by each graph in the ensemble [20–24]; this strong requirement, however, comes at the price of nonergodicity, high computational demand, and poor generalisability [25].

So far, only linear ERGs have been extensively employed to gain insight into the structural organization of real-world systems: For simple graphs, the most important one is the undirected binary configuration model (UBCM), inducing an ensemble of configurations specified by the degree sequence; a useful fitness-induced variant of it, to be employed in presence of partial information is, instead, the density-corrected gravity model (dcGM), solely enforcing (a proxy of) the link density while relying on node-specific quantities identified with the node strengths $\{s_i\}_{i=1}^N$, defined as $s_i = \sum_{j(\neq i)} w_{ij}$, $\forall i$ and representing the total volume of interactions involving each node. This choice is motivated by the evidence that in economic and financial applications a node strength is observable from balance-sheets, transactional or input-output records, while the number of its counterparties (i.e., its degree) is typically not disclosed [26–31]: Treating strengths as exogenous fitnesses, thus, allows one to account for the heterogeneity of nodes without imposing any local constraint [19]: In formulas,

$$p_{ij}^{\text{dcGM}} = \frac{z s_i s_j}{1 + z s_i s_j}, \quad (1)$$

where z can be determined by requiring $L = \sum_i \sum_{j(\neq i)} p_{ij}^{\text{dcGM}} = \langle L \rangle$ [27,32].

Although powerful enough to reproduce many quantities of interest as accurately as the UBCM, the dcGM fails in reproducing the variance of the empirical degree distribution, either overestimating or underestimating it (see Fig. 1). In the first case, larger degrees are overestimated while smaller degrees are underestimated; in the second one, larger degrees

*Contact author: mattia.marzi@imtlucca.it

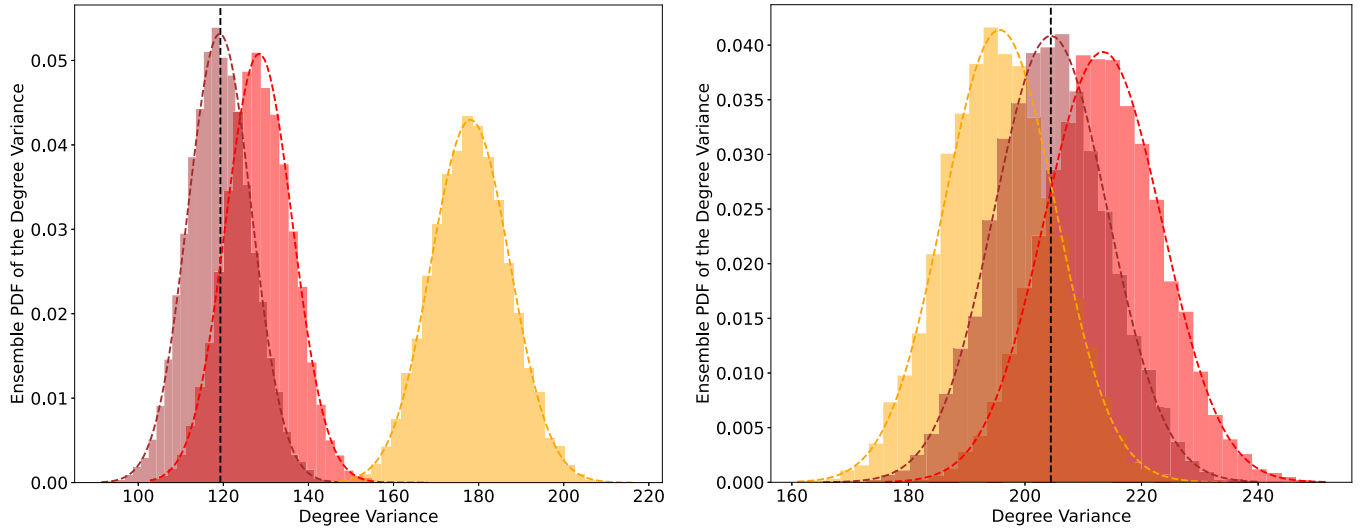


FIG. 1. Analysis of eMID during the 31st week (left panel) and the 42nd week (right panel) of the year 2004. Graphical representation of the agreement between the sample variance $\text{Var}[\mathbf{k}]$ (black, dashed, vertical line) and its expected value $\langle \text{Var}[\mathbf{k}] \rangle$ under the UBCM (red), the dcGM (yellow), and the fit2SM (brown): while the UBCM steadily overestimates it and the dcGM either overestimates or underestimates it, the fit2SM correctly reproduces it. Each ensemble distribution is well approximated by a Gaussian whose parameters match the corresponding average and standard deviation.

are underestimated while smaller degrees are overestimated. Accurately reproducing the variance of the degree distribution is crucial for a variety of applications, such as correctly assessing the systemic risk related to the interconnections of a given financial system [7,33] or the epidemic threshold associated with the interlinkages of a given social system—reading $\langle k \rangle / \langle k^2 \rangle$, hence being governed by the first two moments of the degree distribution [4,34]; another example is provided by the expected consensus time of the so-called coalescing random walks, whose estimation requires calculating the first two moments of the degree distribution as well [35,36].

The aim of this work is to investigate whether it is possible to define a minimal model capable of accurately reproducing both the first and second moments of empirical degree distributions, without resorting to microcanonical constraints. To this end, we consider the class of ERGs with nonlinear constraints and introduce a fitness-induced variant of the two-star model that can be efficiently and reliably implemented. We test it on the transaction-level data constituting the overnight segment of the Electronic Market for Interbank Deposits (eMID), a screen-based market for unsecured deposits [8,37,38]. Although the raw records are directed and weighted (lender, borrower, notional amount), here, we symmetrize and binarize exposures within each, considered time window, i.e., “daily,” “weekly,” “monthly,” “quarterly,” and “yearly” (the description of each aggregation procedure is reported in Appendix A).

II. THE SIMPLEST NONLINEAR CONSTRAINT

The simplest nonlinear constraint is represented by the total number of two-stars [10,39], i.e.,

$$S = \sum_i \sum_{j(>i)} V_{ij} = \sum_i \sum_{j(>i)} \sum_m a_{im} a_{jm}; \quad (2)$$

since

$$\begin{aligned} S &= \frac{1}{2} \sum_m \sum_i \sum_{j(\neq i)} a_{im} a_{jm} \\ &= \frac{1}{2} \sum_m \sum_i a_{im} \left(\sum_j a_{jm} - a_{im} \right) \\ &= \frac{1}{2} \sum_m \left(\sum_i a_{im} k_m - \sum_i a_{im} \right) \\ &= \frac{1}{2} \sum_m (k_m^2 - k_m) \\ &= \frac{1}{2} \sum_m k_m^2 - L, \end{aligned} \quad (3)$$

inverting such a relationship leads the second moment of the empirical degree distribution to be rewritable as

$$\overline{k^2} = \frac{\sum_m k_m^2}{N} = \frac{2S}{N} + \frac{2L}{N} \quad (4)$$

and its variance (hereby, sample variance) as

$$\text{Var}[\mathbf{k}] = \overline{k^2} - \bar{k}^2 = \frac{2S}{N} + \frac{2L}{N} \left(1 - \frac{2L}{N} \right). \quad (5)$$

As a consequence, its expected value reads

$$\begin{aligned} \langle \text{Var}[\mathbf{k}] \rangle &= \frac{2\langle S \rangle}{N} + \frac{2\langle L \rangle}{N} - \frac{4\langle L^2 \rangle}{N^2} \\ &= \frac{2\langle S \rangle}{N} + \frac{2\langle L \rangle}{N} - \frac{4\text{Var}[L]}{N^2} - \frac{4\langle L \rangle^2}{N^2} \\ &= \frac{2\langle S \rangle}{N} - \frac{4\text{Var}[L]}{N^2} + \frac{2\langle L \rangle}{N} \left(1 - \frac{2\langle L \rangle}{N} \right), \end{aligned} \quad (6)$$

an expression showing that reproducing the sample variance requires (at least) the total number of links to be reproduced *exactly*.

III. A MICROCANONICAL APPROACH

Since the result above suggests the microcanonical road as the only viable one, a natural choice would be that of considering the microcanonical version of the two-star model (m2SM), constraining both L and S exactly: Since $\langle L \rangle = L$, $\text{Var}[L] = 0$, and $\langle S \rangle = S$, one would obtain $\langle \text{Var}[\mathbf{k}] \rangle = \text{Var}[\mathbf{k}]$. This model is, however, *homogeneous* and, therefore, unable to reproduce the degree sequence.

We would be, thus, tempted to refine it by constraining *both* the total number of two-stars *and* the degree sequence. Yet, since

$$\begin{aligned} S - \langle S \rangle &= \frac{1}{2} \left[\sum_i (k_i^2 - k_i) - \sum_i (\langle k_i^2 \rangle - \langle k_i \rangle) \right] \\ &= \frac{1}{2} \sum_i (k_i^2 - \langle k_i^2 \rangle) \\ &= \frac{1}{2} \sum_i (\langle k_i \rangle^2 - \langle k_i^2 \rangle) \\ &= -\frac{1}{2} \sum_i \text{Var}[k_i], \end{aligned} \quad (7)$$

satisfying both (sets of) constraints microcanonically is equivalent to requiring that $\text{Var}[k_i] = 0$, $\forall i$. In words, our results indicate that the microcanonical version of the configuration model (mCM) is the (simplest) one guaranteeing $\langle k_i \rangle = k_i$, $\forall i$ and $\langle S \rangle = S$. Moreover, as the relationship

$$\text{Var}[L] = \text{Var} \left[\frac{\sum_i k_i}{2} \right] = \frac{\sum_i \text{Var}[k_i] + 2 \sum_{l < n} \text{Cov}[k_l, k_n]}{4} \quad (8)$$

leads to the expression

$$\begin{aligned} \langle \text{Var}[\mathbf{k}] \rangle &= \frac{2\langle S \rangle}{N} - \frac{\sum_i \text{Var}[k_i] + 2 \sum_{l < n} \text{Cov}[k_l, k_n]}{N^2} \\ &\quad + \frac{2\langle L \rangle}{N} \left(1 - \frac{2\langle L \rangle}{N} \right), \end{aligned} \quad (9)$$

the mCM would also guarantee that $\langle \text{Var}[\mathbf{k}] \rangle = \text{Var}[\mathbf{k}]$: In words, constraining the entire degree sequence *exactly* would lead to reproduce the sample variance as well.

Although the (microcanonical) requirements are clear, the way to realize them is much less so; moreover, the problems affecting microcanonical approaches mentioned in the introductory paragraph lead us to discard this route. Is a canonical way out viable?

IV. THE MEAN-FIELD APPROXIMATION

Handling nonlinear constraints within a canonical framework is usually achieved by adopting the mean-field approximation (see also Appendix B) [10,39].

Constraining both the number of two-stars and the degree sequence within such a framework leads to the degree-corrected two-star model (dc2SM): Formally introduced in Appendix C, its generic probability coefficient reads

$$p_{ij}^{\text{dc2SM}} = \frac{e^{-(\alpha_i + \alpha_j) - \beta(k_i + k_j)}}{1 + e^{-(\alpha_i + \alpha_j) - \beta(k_i + k_j)}} = \frac{x_i x_j y^{k_i + k_j}}{1 + x_i x_j y^{k_i + k_j}}; \quad (10)$$

the dc2SM is, however, unable to match both (sets of) constraints, as reproducing the empirical number of two-stars implies underestimating the degree of at least one node (see also Appendix C). The explanation of such a behavior lies in a simple observation: As the mean-field approximation scheme requires that

$$\text{Var}[k_i] = \sum_{j(\neq i)} p_{ij}(1 - p_{ij}), \quad \forall i, \quad (11)$$

letting Eq. (7) vanish requires the generic probability coefficient of a canonical model to degenerate, becoming either 0 or 1. In other words, constraining both the degrees and the number of two-stars within the mean-field approximation scheme forces the model to become deterministic, individuating the observed configuration [40] $\mathbf{A}^*$ as the only admissible one.

V. A CANONICAL WAY OUT

Let us now ask ourselves if an alternative, canonical road exists. To this aim, let us first notice that $\text{Var}[L]$ is divided by N^2 in Eq. (6): such an addendum may, thus, be expected not to play a relevant role. As a consequence, we are allowed to consider the expression:

$$\langle \text{Var}[\mathbf{k}] \rangle \simeq \frac{2\langle S \rangle}{N} + \frac{2\langle L \rangle}{N} \left(1 - \frac{2\langle L \rangle}{N} \right), \quad (12)$$

requiring the total number of links and the total number of two-stars to be reproduced *on average*. Since any linear ERG reproduces the total number of links, the failure in reproducing the sample variance boils down to the failure in reproducing the total number of two-stars: An overestimation/underestimation of the latter leads, in fact, to an overestimation/underestimation of the former—this is precisely the case of the canonical version of the configuration model (UCBM), that overestimates the total number of two-stars [see Eq. (7)], hence predicting $\langle \text{Var}[\mathbf{k}] \rangle \geq \text{Var}[\mathbf{k}]$ (see Fig. 1).

At this point, a natural choice would be that of considering the canonical version of the two-star model (2SM): Within the mean-field approximation scheme, however, the 2SM is defined by the position

$$p = \frac{xy^{k_i + k_j}}{1 + xy^{k_i + k_j}} = \frac{xy^{2(N-1)p}}{1 + xy^{2(N-1)p}}, \quad (13)$$

which, again, makes it a homogeneous model [41] (see also Appendix C) [10,39].

Constructing a fitness-induced variant of the 2SM (fit2SM) is, nevertheless, feasible: To this aim, it is enough to replace x_i with $\sqrt{z} s_i$ [42] and k_i with its expectation under the same model in Eq. (10). These substitutions lead to the generic

probability coefficient reading

$$p_{ij}^{\text{fit2SM}} = \frac{z s_i s_j y^{\kappa_i + \kappa_j}}{1 + z s_i s_j y^{\kappa_i + \kappa_j}}, \quad (14)$$

where $\langle k_i \rangle_{\text{fit2SM}} \equiv \kappa_i$. The two global parameters z and y can be estimated by solving the (system constituted by the) two nonlinear, coupled equations reading

$$\begin{aligned} L^* &= \sum_i \sum_{j(>i)} p_{ij}^{\text{fit2SM}} = \langle L \rangle, \\ S^* &= \sum_i \sum_{j(>i)} \sum_m p_{im}^{\text{fit2SM}} p_{jm}^{\text{fit2SM}} = \langle S \rangle, \end{aligned} \quad (15)$$

and rewritable [43] as

$$\begin{aligned} z_{(t+1)} &= \frac{L^*}{\sum_i \sum_{j(>i)} \frac{s_i^* s_j^* y_{(t)}^{\kappa_i + \kappa_j}}{1 + z_{(t)} s_i^* s_j^* y_{(t)}^{\kappa_i + \kappa_j}}}, \\ y_{(t+1)} &= \frac{S^*}{\sum_i \sum_{j(>i)} \sum_m \frac{z_{(t)}^2 (s_m^*)^2 y_{(t)}^{\kappa_i + 2\kappa_m + \kappa_j - 1}}{(1 + z_{(t)} s_i^* s_m^* y_{(t)}^{\kappa_i + \kappa_m})(1 + z_{(t)} s_j^* s_m^* y_{(t)}^{\kappa_j + \kappa_m})}}, \end{aligned} \quad (16)$$

Notice that a third set of consistency equations concerning the degrees should be added. Operatively, we should (1) initialize the degrees in some way; (2) solve the system above, obtaining z and y at the “epoch” 1 (i.e., z_1 and y_1); (3) repeat the steps (1) and (2) to calculate the corresponding values of the degrees and use them to obtain z and y at the subsequent “epochs.” In formulas,

$$\kappa_i^{(e+1)} = \sum_{j(\neq i)} \frac{z_{(e)} s_i^* s_j^* y_{(e)}^{\kappa_i^{(e)} + \kappa_j^{(e)}}}{1 + z_{(e)} s_i^* s_j^* y_{(e)}^{\kappa_i^{(e)} + \kappa_j^{(e)}}}, \quad \forall i, \quad (17)$$

with e indicating the “epoch” of the resolution. As we show in Appendix D, however, stopping at the “epoch” 0, by posing

$$\kappa_i^{(0)} = \sum_{j(\neq i)} p_{ij}^{\text{dcGM}}, \quad \forall i \quad (18)$$

is (already) enough to achieve good results in terms of speed and accuracy of resolution. The system above ensures that the fit2SM correctly reproduces the total number of links, the total number of two-stars and the sample variance, besides accounting for the nodes’ heterogeneity.

VI. TESTING THE PERFORMANCE OF THE FIT2SM

A. Reproducing a network’s local properties

Since the degree sequence is a standard diagnostic in network reconstruction exercises, we start our analysis by testing the accuracy of the fit2SM in reproducing it. We do so by means of the quantile-quantile (QQ) plot: given the empirical sequence $\{k_i\}_{i=1}^N$ and a model-based, expected one $\{\langle k_i \rangle\}_{i=1}^N$, we sort both in increasing order and scatter the corresponding pairs of values [44]; while perfect agreement would place all points on the identity line, misestimations of the values would be revealed by systematic deviations.

Let us now select five snapshots by drawing one calendar year at random and, conditionally on it, drawing one

quarter, one month, one ISO week, and one trading day at random as well. Figure 2 reports the resulting QQ plots for year 2001, quarter Q3, month April, ISO week 2, and day 2001-11-21; the visual inspection is accompanied by the accuracy indicators named average relative error, defined as $\text{ARE} = N^{-1} \sum_i |k_i - \langle k_i \rangle| / k_i$, and maximum relative error, defined as $\text{MRE} = \max_i \{|k_i - \langle k_i \rangle| / k_i\}$.

While the UBCM reproduces the degrees by construction, the dcGM and the fit2SM yield comparable errors, with the ARE ranging from $\simeq 0.51$ to $\simeq 0.80$ and the MRE from $\simeq 6.60$ to $\simeq 22.2$, across aggregations: more specifically, while the fit2SM achieves a smaller ARE at the daily, monthly, and quarterly scales, the dcGM performs better at the weekly and yearly scales; at the node level, the fit2SM yields a smaller ARE in the 65% of cases at the daily scale, in the 61% of cases at the weekly scale, in the 41% of cases at the monthly scale, in the 38% of cases at the quarterly scale, and in the 54% of cases at the yearly scale. Repeating the procedure above for other, randomly selected snapshots yields analogous results.

As these observations point out, the predictions of the degrees returned by the dcGM and the fit2SM turn out to be very close: as the fit2SM is not conceived to improve the matching of the degrees but to incorporate a nonlinear constraint to reproduce their variance, the related “benefits” are expected to emerge more clearly when inspecting properties that depend explicitly on the second moment—such as the spectral radius (see below).

B. Reproducing a network spectral properties

As a second test, let us inspect the accuracy of the fit2SM in reproducing the spectral radius, hereby indicated with π_1 , of a given network. To this aim, let us follow Ref. [8] and consider the fully controllable case represented by the Chung-Lu model [45]. A way to identify π_1 in case $p_{ij} = k_i k_j / 2L$, $\forall i, j$ rests upon the relationship

$$\mathbf{P} = \frac{\mathbf{k} \otimes \mathbf{k}}{2L} = \frac{|k\rangle\langle k|}{2L}, \quad (19)$$

indicating that, in such a case, the matrix $\mathbf{P}$ can be obtained as the direct product of the vector of degrees with itself. Employing the bra-ket formalism allows the calculations to be carried out quite easily, i.e., as

$$\mathbf{P}|k\rangle = \frac{|k\rangle\langle k|}{2L}|k\rangle = \frac{\langle k|k\rangle}{2L}|k\rangle, \quad (20)$$

where $\langle k|k\rangle = \sum_i k_i^2$. Since $\mathbf{P}$ obeys the Perron-Frobenius theorem [46], the equation above allows us to identify the value of its spectral radius quite straightforwardly as $\pi_1 = \langle k|k\rangle / 2L = \sum_i k_i^2 / 2L$. The Chung-Lu model is, however, defined by the position $p_{ij} = k_i k_j / 2L$, $\forall i \neq j$, a piece of evidence leading us to write

$$\pi_1 \simeq \frac{\langle k|k\rangle}{2L} = \sum_i \frac{k_i^2}{2L} = \frac{N}{2L} \sum_i \frac{k_i^2}{N} = \frac{\overline{k^2}}{\overline{k}}; \quad (21)$$

in other words, the Chung-Lu model represents a convenient benchmark to make the connection between the accuracy of a

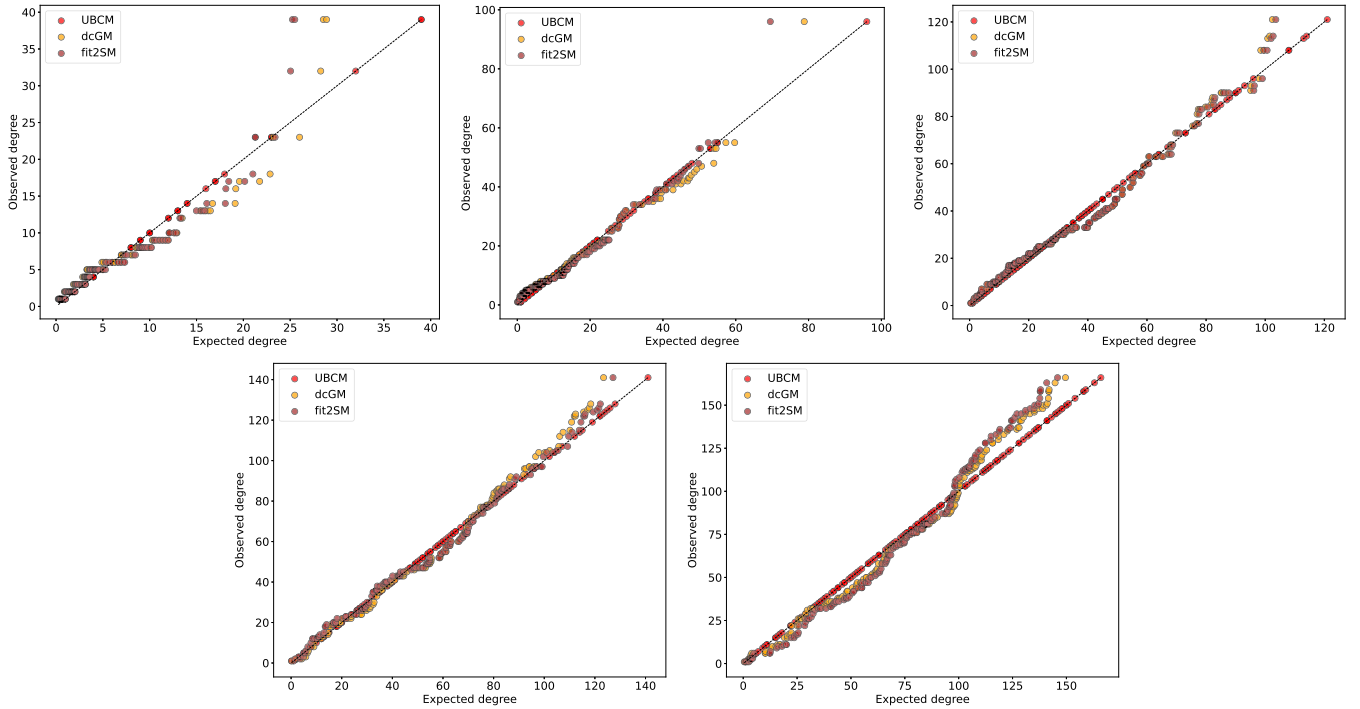


FIG. 2. Quantile-quantile plots comparing the observed degrees and the expected ones under the UBCM (red), the dcGM (yellow), and the fit2SM (brown), across five aggregations of the eMID interbank market (from top-left to bottom-right: daily, weekly, monthly, quarterly, and yearly). These snapshots are selected by drawing one calendar year at random (here, 2001) and, conditionally on it, drawing at random one quarter (here, Q3), one month (here, April), one ISO week (here, week 2), and one trading day (here, 2001-11-21). As it can be appreciated, the dcGM and the fit2SM perform very similarly, and both satisfactorily, in reproducing the degree sequence induced by the corresponding snapshot.

network model in reproducing the first two moments of the degree distribution and the one in reproducing the spectral properties of the corresponding configuration explicit, via the dependence of π_1 on $\bar{k}$ and $\bar{k}^2$.

Let us now consider the eMID snapshot corresponding to the day 2010-07-19 and calibrate the model above on it—a choice allowing us to deal with a graphical degree sequence. Afterwards, let us generate $M = 10^3$ configurations and treat each of them as a plausible, empirical one; solving our models on the latter leads to three distributions of $M = 10^3$ estimates of the spectral radius each (hereby indicated with λ_1 , to distinguish them from the reference value π_1 ; see also Appendix E): as the left panel of Fig. 3 shows, while the UBCM overestimates the spectral radius of the generative model ($\lambda_1^{\text{UBCM}} \simeq 9.80 > \pi_1 \simeq 9.35$) and the dcGM underestimates it ($\lambda_1^{\text{dcGM}} \simeq 8.66 < \pi_1 \simeq 9.35$), the fit2SM correctly reproduces it ($\lambda_1^{\text{fit2SM}} \simeq 9.40 \gtrsim \pi_1 \simeq 9.35$).

A more explicit test concerns the accuracy of the fit2SM in reproducing the empirical spectral radius of the binary, undirected version of eMID across the years 1999–2012: as the system we are considering cannot be expected to perfectly align with the Chung-Lu model, deviations from the previous results are expected; as the calculation of the mean and standard deviation of the absolute error $|\lambda_1^{\text{model}} - \lambda_1^{\text{obs}}|$ over the $M = 10^3$ configurations drawn from the snapshot-specific ensemble of each model reveals, in fact, one finds that

at the daily timescale, the fit2SM achieves the smallest average error (i.e., 0.79 ± 0.51), performing better than the

UBCM and the dcGM in 94.00% and 97.60% of the snapshots, respectively;

at the weekly timescale, the fit2SM achieves the smallest average error (i.e., 0.97 ± 0.66), performing better than the UBCM and the dcGM in 82.40% and 98.70% of the snapshots, respectively;

at the monthly timescale, the fit2SM achieves the smallest average error (i.e., 0.71 ± 0.66), performing better than the UBCM and the dcGM in 64.80% and 74.50% of the snapshots, respectively;

at the quarterly timescale, the UBCM achieves the smallest average error (i.e., 0.43 ± 0.21), the fit2SM performing better than the dcGM in 67.30% of the snapshots but performing better than the UBCM only in 32.70% of the snapshots;

at the yearly timescale, the UBCM achieves the smallest average error (i.e., 0.61 ± 0.57), the fit2SM performing better than the dcGM in 100.00% of the snapshots but performing better than the UBCM only in 14.30% of the snapshots.

A related comparison is depicted in the right panel of Fig. 3, illustrating the distribution—summarized by the violin plot—of the relative error $(\lambda_1^{\text{model}} - \lambda_1^{\text{obs}})/\lambda_1^{\text{obs}}$ in reproducing the empirical spectral radius, across snapshots and temporal aggregations. It is worth noticing that the occasional, superior performance of the UBCM on the densest snapshots comes at the cost of taking the entire degree sequence as input; the fit2SM, on the other hand, solely relies upon two global parameters, a feature making such a model substantially less demanding, in terms of informational requirements and numerical complexity.

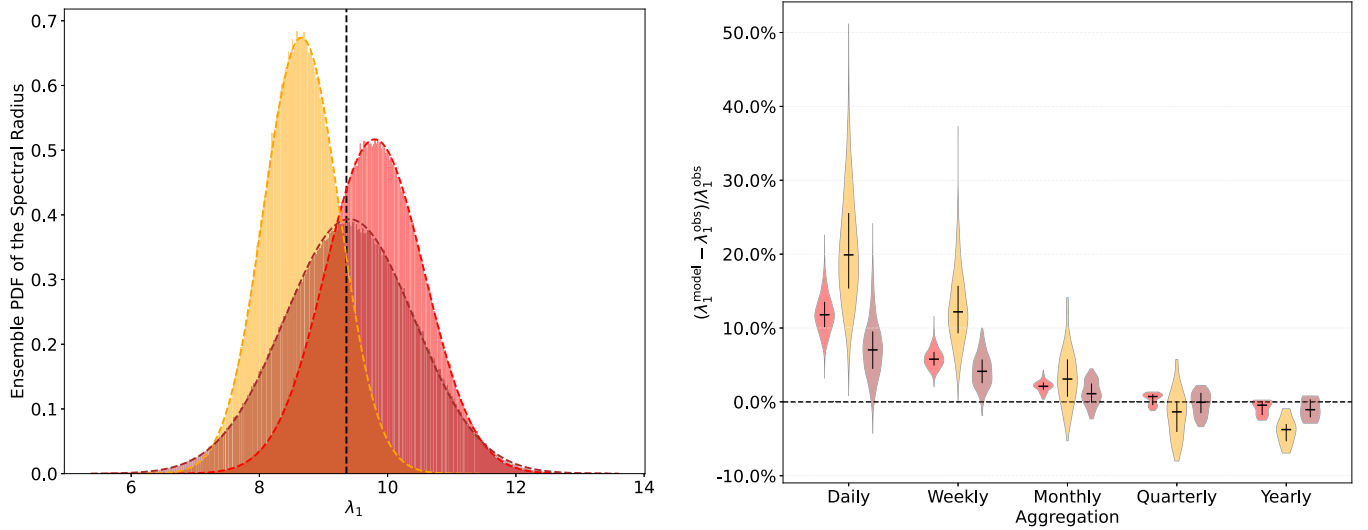


FIG. 3. Left panel: distributions of the spectral radius induced by $M = 10^3$ configurations sampled from the Chung-Lu model defined by the topology of the eMID snapshot corresponding to the day 2010-07-19, according to the UBCM (red), the dcGM (yellow), and the fit2SM (brown). The three models were solved for each of the $M = 10^3$, generated configurations and the corresponding ensembles explicitly sampled in order to obtain an estimation of π_1 : what we find is that $\lambda_1^{\text{dcGM}} < \pi_1 \lesssim \lambda_1^{\text{fit2SM}} < \lambda_1^{\text{UBCM}}$. In words, while the UBCM overestimates the spectral radius of the generative model (black, dashed, vertical line) and the dcGM underestimates it, the fit2SM correctly reproduces it. Each ensemble distribution is well approximated by a Gaussian whose parameters match the corresponding average and standard deviation. Right panel: violin plots (summarizing the distributions) of the relative error in reproducing the empirical spectral radius λ_1^{obs} , across snapshots and temporal aggregations. The fit2SM (brown) yields smaller errors than (1) the dcGM (yellow) at all timescales and (2) the UBCM (red) at the daily and weekly timescales.

C. Reproducing a network structure

A compact indicator of a model's performance in reconstructing a network structure is provided by the Bayesian information criterion (BIC), reading

$$\text{BIC} = c \ln V - 2\mathcal{L}, \quad (22)$$

where $V = N(N-1)/2$ accounts for the system size, the number of parameters amounts at $c = N$ for the UBCM, $c = 1$ for the dcGM, and $c = 2$ for the fit2SM, and the log-likelihood reads

$$\mathcal{L} = \ln P(\mathbf{A}) = \sum_i \sum_{j(>i)} [a_{ij} \ln p_{ij} + (1 - a_{ij}) \ln(1 - p_{ij})] \quad (23)$$

for all of them, with

$$p_{ij}^{\text{UBCM}} = \frac{x_i x_j}{1 + x_i x_j}, \quad (24)$$

$$p_{ij}^{\text{dcGM}} = \frac{z_i s_j}{1 + z_i s_j}, \quad (25)$$

and

$$p_{ij}^{\text{fit2SM}} = \frac{z_i s_j y_i^{\kappa_i^{(0)} + \kappa_j^{(0)}}}{1 + z_i s_j y_i^{\kappa_i^{(0)} + \kappa_j^{(0)}}}. \quad (26)$$

Testing the three models above on eMID reveals systematic, aggregation-dependent differences. To present the full batch of results in a compact and directly comparable way,

Fig. 4 illustrates the distributions of

$$\delta \text{BIC}_{\text{model}} = \frac{\text{BIC}_{\text{model}} - \text{BIC}_{\text{min}}}{\text{BIC}_{\text{model}}}, \quad (27)$$

where $\text{BIC}_{\text{min}} = \min_{\text{model}} \{\text{BIC}_{\text{model}}\}$; while a null value of such an index identifies the best-performing model on the considered snapshot, the larger its value, the worse the performance of the corresponding model. As Fig. 4 shows, our canonical, nonlinear model

steadily outperforms both the dcGM and the (much) more demanding UBCM at the daily and weekly timescales, i.e., on sparser configurations (in fact, $\text{BIC}_{\text{fit2SM}} < \text{BIC}_{\text{dcGM}}$ on 90.76% of daily snapshots and 97.63% of weekly snapshots; $\text{BIC}_{\text{fit2SM}} < \text{BIC}_{\text{UBCM}}$ on 99.97% of daily snapshots and 100.00% of weekly snapshots);

competes with the dcGM at the monthly timescale (in fact, $\text{BIC}_{\text{fit2SM}} < \text{BIC}_{\text{dcGM}}$ on 64.24% of monthly snapshots and $\text{BIC}_{\text{fit2SM}} < \text{BIC}_{\text{UBCM}}$ on 33.94% of monthly snapshots but $\text{BIC}_{\text{fit2SM}} < \text{BIC}_{\text{dcGM}}$ on 100.00% of monthly snapshots and $\text{BIC}_{\text{fit2SM}} < \text{BIC}_{\text{UBCM}}$ on 97.78% of monthly snapshots since 2009);

complete with the dcGM at the quarterly timescale during the last snapshots of our dataset (in fact, $\text{BIC}_{\text{fit2SM}} < \text{BIC}_{\text{dcGM}}$ on 55.56% of quarterly snapshots but $\text{BIC}_{\text{fit2SM}} < \text{BIC}_{\text{UBCM}}$ on 11.11% of quarterly snapshots since 2009).

Overall, these findings suggest the presence of a dependency between the performance of the fit2SM and the network link density, i.e., the sparser the network, the better its performance. To test such an hypothesis, let us scatter the index of

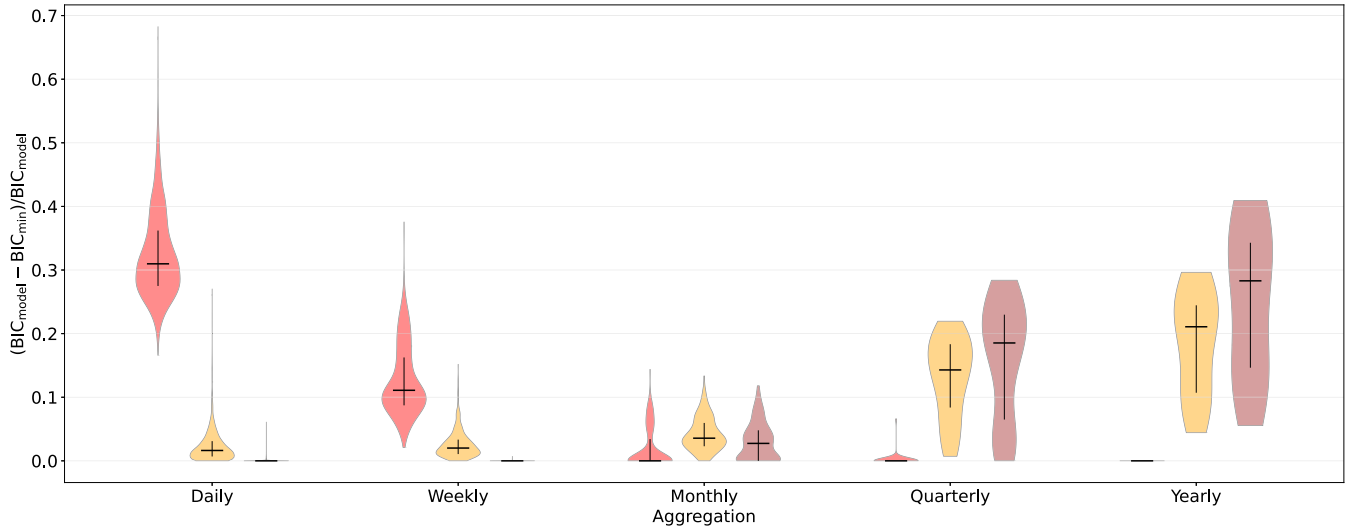


FIG. 4. For each snapshot and aggregation level, we show the violin plots (summarizing the distributions) of the relative error $\delta \text{BIC}_{\text{model}}$ [defined in Eq. (27)]. Upon remembering that a null value of such an index identifies the best-performing model on the considered snapshot, let us notice that the score induced by the fit2SM (brown) concentrates near zero—hence such a model dominates over the competing ones—at finer aggregations, while the UBCM (red) is the winner at coarser aggregations.

relative performance

$$\Delta \text{BIC}_{\text{model}} = \frac{\text{BIC}_{\text{model}} - \text{BIC}_{\text{fit2SM}}}{\text{BIC}_{\text{model}}} \quad (28)$$

versus the relative error in reproducing the total number of two-stars,

$$\Delta S_{\text{model}} = \frac{\langle S \rangle_{\text{model}} - \langle S \rangle_{\text{fit2SM}}}{\langle S \rangle_{\text{model}}} = \frac{\langle S \rangle_{\text{model}} - S^*}{\langle S \rangle_{\text{model}}} \quad (29)$$

for both the UBCM and the dcGM. As Fig. 5 reveals, whenever ΔS_{model} is large—typically the case for sparser configurations—model selection favours the fit2SM.

D. The fit2SM as a generative model

On the basis of the aforementioned results, let us employ the fit2SM as a generative model and test the accuracy of the dcGM in recovering the early-warning signals (EWS) shown in Ref. [8]: To this aim, let us compute the z score reading $z = (\mu - \lambda_1^{\text{fit2SM}})/\sigma$, where μ and σ are the expected value and standard deviation of the dcGM estimate computed over the

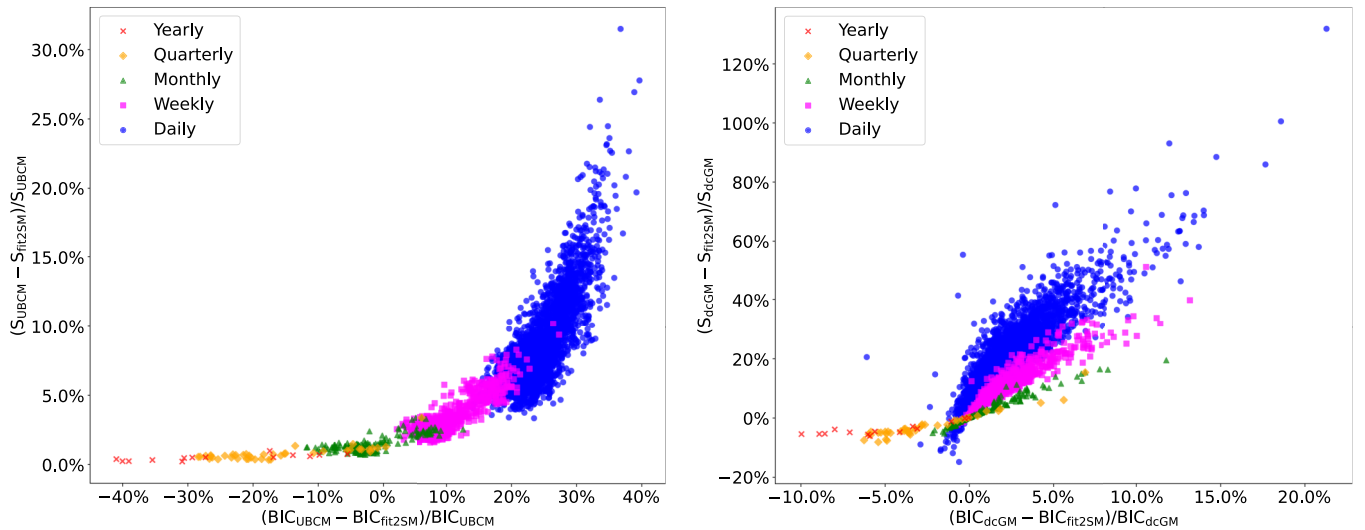


FIG. 5. Left panel: relative improvement of the fit2SM vs the relative error in estimating the number of two-stars with respect to the UBCM. Right panel: relative improvement of the fit2SM vs the relative error in estimating the number of two-stars with respect to the dcGM. A trend appear: the larger the error in estimating the total number of two-stars, the more the fit2SM outperforms the other models, an evidence suggesting that a poorer estimation of the total number of two-stars corresponds to a worse statistical fit of the eMID structure. This effect is more pronounced for sparser configurations, on which both the UBCM and the dcGM tend to commit larger errors—while the UBCM steadily overestimates S , the dcGM either overestimate or underestimate it. Each point is a network at a given level of aggregation.

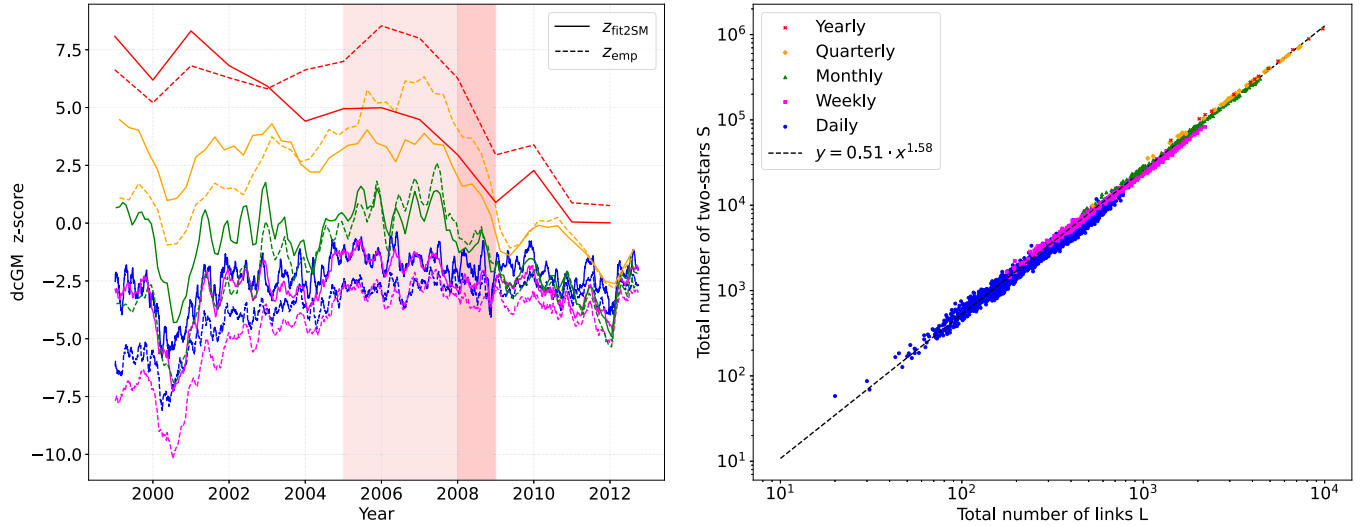


FIG. 6. Left panel: solid lines represent the z scores reading $z = (\mu - \lambda_1^{\text{fit2SM}})/\sigma$, while dashed lines represent the z scores reading $z = (\mu - \lambda_1)/\sigma$. In both cases, μ and σ have been calculated via the dcGM. The similarity of the two series of values indicate that the fit2SM represents a valid generative model (e.g., for early-warning signals detection) in case the true network topology is unavailable. The precrisis period 2005–2008 and the crisis period 2008–2009 are highlighted by shaded areas. Different colours indicate different timescale: blue-daily, magenta-weekly, green-monthly, orange-quarterly, red-yearly. Trends have been smoothed via a rolling average over the points $[t - 10, t + 10]$ for the daily aggregation, over the points $[t - 7, t + 7]$ for the weekly aggregation, over the points $[t - 5, t + 5]$ for the monthly aggregation, over the points $[t - 3, t + 3]$ for the quarterly aggregation. Right panel: irrespectively from the aggregation level at which eMID is considered, a relationship like $S = aL^b$ seems to hold true, the average parameters reading $a \simeq 0.51$ and $b \simeq 1.58$. Each point is a network at a given level of aggregation.

eMID surrogate generated by the fit2SM. As the left panel of Fig. 6 shows, the z score induced by the fit2SM closely mirrors the one evaluated with respect to the true network topology for all timescale, a result confirming the robustness of the fit2SM as a reference model for EWS detection. Stated otherwise, in case the true network topology is not available, the fit2SM proxies quite effectively the ground truth.

VII. DISCUSSION

After reviewing a number of negative results about the possibility of defining an ERG constraining both the degree sequence and the total number of two-stars, we propose a minimal, canonical model capable of reproducing the variance of a degree distribution while accounting for the nodes heterogeneity: In order to achieve such a goal, a fitness-induced variant of the 2SM, named fit2SM (i.e., a “softened” nonlinear ERG whose definition is inspired to its fitness-based, linear counterpart), has been employed.

Besides achieving a large accuracy in reproducing a network structure at different levels, the fit2SM seems also to mitigate a problem affecting the dcGM, i.e., that of returning sampled configurations with too many isolated nodes: In fact, $\langle N_{\text{fit2SM}}^0 \rangle < \langle N_{\text{dcGM}}^0 \rangle$ on 97.63% of daily snapshots, 98.61% of weekly snapshots, and 67.88% of monthly snapshots; interestingly, $\langle N_{\text{fit2SM}}^0 \rangle < \langle N_{\text{UBCM}}^0 \rangle$ on 58.79% of monthly snapshots and 38.10% of quarterly snapshots but $\langle N_{\text{fit2SM}}^0 \rangle < \langle N_{\text{UBCM}}^0 \rangle$ on 66.67% of monthly snapshots and 55.56% of quarterly snapshots since 2009 (see also Appendix F).

A last comment concerns the availability of the information about the total number of two-stars: this is rarely the case. As the right panel of Fig. 6 shows, however, it can be deduced

from the one concerning the total number of links: irrespectively from the aggregation level, in fact, a relationship like $S = aL^b$ seems to hold true, the average parameters reading $a \simeq 0.51$ and $b \simeq 1.58$ (see also Appendix D). An analogous relationship holds true for the data concerning the yearly snapshots of the International Trade Network from 1990 to 2000 [8], the average parameters, now, reading $a \simeq 0.44$ and $b \simeq 1.61$.

Future research will explore nonlinear models enforcing more complex patterns (e.g., the Strauss one, constraining the total number of triangles).

ACKNOWLEDGMENTS

M.M., D.G., and T.S. acknowledge support from the project “SoBigData.it—Strengthening the Italian RI for Social Mining and Big Data Analytics”—IR0000013—CUP B53C22001760006, financed by European Union—Next Generation EU—National Recovery and Resilience Plan (Piano Nazionale di Ripresa e Resilienza, PNRR)—M4C2 I.3.1; D.G. and T.S. acknowledge support from the projects “Reconstruction, Resilience and Recovery of Socio-Economic Networks” RECON-NET EP_FAIR_005—PE0000013 “FAIR,” Funded by the European Union Next Generation EU, PNRR Mission 4 Component 2, Investment 3.1; “RE-Net: Reconstructing economic networks: From physics to machine learning and back”—2022MTBB22, Funded by the European Union Next Generation EU, PNRR Mission 4 Component 2 Investment 1.1, CUP: D53D23002330006; “C2T—From Crises to Theory: Towards a science of resilience and recovery for economic and financial systems”—P2022E93B8, Funded by the European Union Next Generation EU, PNRR Mission

4 Component 2 Investment 1.1, CUP: D53D23019330001. This work was also supported by the European Union—NextGenerationEU and funded by the Italian Ministry of University and Research (MUR)—National Recovery and Resilience Plan (Piano Nazionale di Ripresa e Resilienza, PNRR), project “A Multiscale Integrated Approach to the Study of the Nervous System in Health and Disease” (MNESYS, PE0000006, DN. 1553, 11/10/2022). We thank Giulio Cimini for fruitful and insightful discussions.

Study conception and design: M.M., F.G., D.G., T.S. Data collection: M.M. Analysis and interpretation of results: M.M., F.G., D.G., T.S. Draft manuscript preparation: M.M., F.G., T.S.

The authors declare no competing interests.

DATA AVAILABILITY

The Python package named `fit2SM`, implementing the algorithms described in the main text, is available on PyPI [47]. The data that support the findings of this article are not publicly available because they are owned by a third party and the terms of use prevent public distribution. The data are available from the authors upon reasonable request.

APPENDIX A: DATA DESCRIPTION AND TEMPORAL AGGREGATIONS

We employ transaction-level records from the eMID), a screen-based market for unsecured deposits [8,37,38], restricting the analysis to the segment that represents the vast majority of the activity on such a market, i.e., the overnight one. Our sample spans trading days from January 1999 to September 2012: each trading day d defines a weighted, directed matrix $\mathbf{V}(d)$ whose entry $v_{ij}(d)$ equals the total notional amount lent by bank i to bank j on day d ; naturally, $v_{ii}(d) = 0$.

Throughout the paper, the labels “daily”, “weekly”, “monthly”, “quarterly” and “yearly” refer to calendar aggregations of daily records: more specifically, “weekly” refers to ISO calendar weeks (Monday to Sunday), “monthly” to calendar months, “quarterly” to standard quarters (Q1: Jan-Mar; Q2: Apr-Jun; Q3: Jul-Sep; Q4: Oct-Dec), and “yearly” to calendar years. Importantly, trading on eMID only occurs on business days, so that weekends and bank holidays do not contribute to observations: consequently, the number of trading days within a given calendar window is not constant across windows of the same kind; in practice, we aggregate over the set of trading days that are present in the raw records during a specific window. Let Δ_t denote the set of trading days belonging to the calendar window indexed by t (week, month, quarter, or year). We, thus, aggregate weights according to

$$v_{ij}^\Delta(t) = \sum_{d \in \Delta_t} v_{ij}(d); \quad (\text{A1})$$

as a last observation, let us stress that, for each window, we restrict the set of nodes to the banks that are active within that window, i.e., that are involved in at least one transaction during Δ_t : the number of nodes, thus, becomes a time-dependent quantity.

Since the empirical records are weighted and directed, we construct an undirected exposure matrix by symmetrizing the aggregated weights as

$$w_{ij}^\Delta(t) = v_{ij}^\Delta(t) + v_{ji}^\Delta(t), \quad i \neq j, \quad (\text{A2})$$

and define the corresponding binary adjacency matrix as

$$a_{ij}^\Delta(t) = \mathbb{1}\{w_{ij}^\Delta(t) > 0\}, \quad i \neq j \quad (\text{A3})$$

with $a_{ii}^\Delta(t) = 0$. Naturally, we compute the node strengths from the underlying, symmetrized weights as

$$s_i^\Delta(t) = \sum_{j(\neq i)} w_{ij}^\Delta(t), \quad (\text{A4})$$

and use them as the exogenous fitnesses informing both the dcGM and the fit2SM.

APPENDIX B: FROM LINEAR HAMILTONIANS TO SEPARABLE HAMILTONIANS

To ease the mathematical manipulations, let us focus on binary, undirected networks. Linear ERGs are described by Hamiltonians reading

$$H(\mathbf{A}) = \sum_i \sum_{j(>i)} H_{ij}(a_{ij}) = \sum_i \sum_{j(>i)} \theta_{ij} a_{ij}, \quad (\text{B1})$$

hence inducing models that can be factorized as

$$P(\mathbf{A}) = \prod_i \prod_{j(>i)} p_{ij}^{a_{ij}} (1 - p_{ij})^{1-a_{ij}}, \quad (\text{B2})$$

with

$$p_{ij} = \frac{e^{-\theta_{ij}}}{1 + e^{-\theta_{ij}}} = \frac{x_{ij}}{1 + x_{ij}}.$$

The related ensembles are easy to sample since $a_{ij} \sim \text{Ber}[p_{ij}]$ and $a_{ij} \perp\!\!\!\perp a_{kl}$, with $(i, j) \neq (k, l)$. Separable Hamiltonians, instead, read

$$H(\mathbf{A}) = \sum_i \sum_{j(>i)} H_{ij}(\mathbf{A}) = \sum_i \sum_{j(>i)} a_{ij} \times f_{ij}(\mathbf{A}, \theta_{ij}), \quad (\text{B3})$$

hence inducing models that can be factorized as well as

$$P(\mathbf{A}) = \prod_i \prod_{j(>i)} p_{ij}^{a_{ij}} (1 - p_{ij})^{1-a_{ij}}, \quad (\text{B4})$$

but with

$$p_{ij} = \frac{e^{-f_{ij}(\mathbf{P}, \theta_{ij})}}{1 + e^{-f_{ij}(\mathbf{P}, \theta_{ij})}}, \quad (\text{B5})$$

where $\mathbf{P} = \langle \mathbf{A} \rangle$. The mean-field approximation consists precisely in replacing $\mathbf{A}$ with $\mathbf{P}$ in f_{ij} in Eq. (B3).

APPENDIX C: THE TWO-STAR MODEL AND ITS DEGREE-CORRECTED VERSION

Let us provide a concrete example of the mean-field approximation above, by considering the so-called two-star

model (2SM). It is defined by

$$\begin{aligned}
 H_{2SM}(\mathbf{A}) &= \theta L + \psi S \\
 &= \theta L + \psi \sum_i \sum_{l(>i)} V_{il} \\
 &= \theta L + \psi \sum_i \sum_{l(>i)} \sum_j a_{ij} a_{jl} \\
 &= \theta L + \frac{\psi}{2} \sum_i \sum_{l(\neq i)} \sum_j a_{ij} a_{jl} \\
 &= \theta L + \frac{\psi}{2} \sum_j \sum_i \sum_{l(\neq i)} a_{ij} a_{jl} \\
 &= \theta L + \frac{\psi}{2} \sum_j \sum_i \left[\sum_l a_{ij} a_{jl} - a_{ji} \right] \\
 &= \theta L + \frac{\psi}{2} \left[\sum_i \sum_j \sum_l a_{ij} a_{jl} - \sum_j k_j \right] \\
 &= \theta L + \frac{\psi}{2} \left[\sum_i \sum_j a_{ij} k_j - \sum_j k_j \right] \\
 &= \theta L + \frac{\psi}{2} \left[\sum_i \sum_{j(>i)} (a_{ij} k_j + a_{ji} k_i) - \sum_j k_j \right] \\
 &= \theta L + \frac{\psi}{2} \left[\sum_i \sum_{j(>i)} a_{ij} (k_i + k_j) - \sum_j k_j \right] \\
 &= \theta L + \frac{\psi}{2} \left[\sum_i \sum_{j(>i)} a_{ij} (k_i + k_j) \right] - \frac{\psi}{2} \sum_j k_j \\
 &= \theta L + \frac{\psi}{2} \left[\sum_i \sum_{j(>i)} a_{ij} (k_i + k_j) \right] - \psi L \\
 &= (\theta - \psi) L + \frac{\psi}{2} \left[\sum_i \sum_{j(>i)} a_{ij} (k_i + k_j) \right], \quad (C1)
 \end{aligned}$$

upon renaming $\theta - \psi$ as α and $\psi/2$ as β , one finds

$$\begin{aligned}
 H_{2SM}(\mathbf{A}) &= \alpha L + \beta \left[\sum_i \sum_{j(>i)} a_{ij} (k_i + k_j) \right] \\
 &= \alpha \sum_i \sum_{j(>i)} a_{ij} + \beta \left[\sum_i \sum_{j(>i)} a_{ij} (k_i + k_j) \right] \\
 &= \sum_i \sum_{j(>i)} a_{ij} [\alpha + \beta(k_i + k_j)], \quad (C2)
 \end{aligned}$$

and identifying f_{ij} with $\alpha + \beta(k_i + k_j)$ leads to

$$p_{ij}^{2SM} = \frac{e^{-\alpha - \beta(k_i + k_j)}}{1 + e^{-\alpha - \beta(k_i + k_j)}} = \frac{xy^{k_i + k_j}}{1 + xy^{k_i + k_j}}. \quad (C3)$$

Since the information about the degrees is not supposed to be available, consistency requires that

$$p = \frac{xy^{2(N-1)p}}{1 + xy^{2(N-1)p}}, \quad (C4)$$

in case the information about the degrees is, instead, available, the Hamiltonian becomes

$$H_{dc2SM}(\mathbf{A}) = \sum_i \theta_i k_i + \psi S \quad (C5)$$

and induces the expression

$$p_{ij}^{dc2SM} = \frac{e^{-(\alpha_i + \alpha_j) - \beta(k_i + k_j)}}{1 + e^{-(\alpha_i + \alpha_j) - \beta(k_i + k_j)}} = \frac{x_i x_j y^{k_i + k_j}}{1 + x_i x_j y^{k_i + k_j}}. \quad (C6)$$

Let us now consider that

$$\begin{aligned}
 \sum_i \binom{\langle k_i \rangle}{2} &= \frac{1}{2} \sum_i [\langle k_i \rangle^2 - \langle k_i \rangle] \\
 &= \frac{1}{2} \sum_i [\langle k_i^2 \rangle - \text{Var}[k_i] - \langle k_i \rangle] \\
 &= \langle S \rangle - \frac{1}{2} \sum_i \text{Var}[k_i]. \quad (C7)
 \end{aligned}$$

As a consequence, the expected number of two-stars reads

$$\langle S \rangle = \sum_i \binom{\langle k_i \rangle}{2} + \frac{1}{2} \sum_i \text{Var}[k_i]. \quad (C8)$$

Since $S = \sum_i \binom{k_i}{2}$, requiring $\langle k_i \rangle = k_i$ leads one to obtain

$$\langle S \rangle = \sum_i \binom{k_i}{2} + \frac{1}{2} \sum_i \text{Var}[k_i] = S + \frac{1}{2} \sum_i \text{Var}[k_i] \geq S, \quad (C9)$$

the equivalence holding true in case $\text{Var}[k_i] = 0$, $\forall i$ (i.e., either in the microcanonical case or in the canonical, deterministic case). If, instead, the following relationship holds true:

$$\sum_i \binom{\langle k_i \rangle}{2} < \sum_i \binom{k_i}{2}, \quad (C10)$$

one finds that

$$\begin{aligned}
 \langle S \rangle &= \sum_i \binom{\langle k_i \rangle}{2} + \frac{1}{2} \sum_i \text{Var}[k_i] \\
 &< \sum_i \binom{k_i}{2} + \frac{1}{2} \sum_i \text{Var}[k_i] = S + \frac{1}{2} \sum_i \text{Var}[k_i], \quad (C11)
 \end{aligned}$$

i.e.,

$$\langle S \rangle < S + \frac{1}{2} \sum_i \text{Var}[k_i], \quad (C12)$$

or, even more explicitly,

$$\langle S \rangle - S < \frac{1}{2} \sum_i \text{Var}[k_i]. \quad (C13)$$

In order for Eq. (C10) to hold true, the condition $\langle k_i \rangle < k_i$ must be verified for *at least* one node: In other words, reproducing S requires *at least* one degree to be underestimated.

TABLE I. Performance of the fixed-point algorithm to solve the systems of equations defining the fit2SM on several snapshots of eMID (N is the total number of nodes, L is the total number of links, and S is the total number of two-stars).

	N	L	$\langle L \rangle$	MRE_L	S	$\langle S \rangle$	MRE_S	Time (s)
eMID 1999—full year	215	9770	9770	9.08×10^{-11}	1168 583	1168 583	4.95×10^{-11}	5.69
eMID 1999—2nd quarter	200	6536	6536	3.43×10^{-10}	598 078	598 078	1.42×10^{-10}	2.23
eMID 1999—5th month	194	4153	4153	7.55×10^{-10}	271 858	271 858	1.96×10^{-10}	0.65
eMID 1999—15th week	191	1920	1920	1.42×10^{-09}	67 109	67 109	9.85×10^{-11}	0.41
eMID 1999—160th day	174	629	629	7.78×10^{-10}	10 537	10 537	4.57×10^{-11}	0.16
eMID 2004—full year	175	4327	4327	2.41×10^{-10}	325 442	325 442	1.17×10^{-10}	4.63
eMID 2004—2nd quarter	163	2948	2948	5.22×10^{-10}	170 352	170 352	1.99×10^{-10}	0.79
eMID 2004—5th month	152	1995	1995	1.01×10^{-09}	85 170	85 170	1.85×10^{-10}	0.36
eMID 2004—15th week	135	925	925	1.70×10^{-09}	21 199	21 199	9.68×10^{-11}	0.15
eMID 2004—160th day	120	407	407	1.85×10^{-09}	5014	5014	4.82×10^{-10}	0.10
eMID 2012—full year	97	1421	1421	5.85×10^{-10}	57 671	57 671	1.56×10^{-10}	0.41
eMID 2012—2nd quarter	85	912	912	5.70×10^{-10}	27 417	27 417	5.62×10^{-11}	0.19
eMID 2012—5th month	82	545	545	6.58×10^{-10}	10 474	10 474	9.89×10^{-11}	0.12
eMID 2012—15th week	66	238	238	1.14×10^{-09}	2859	2859	1.28×10^{-10}	0.07
eMID 2012—160th day	54	103	103	2.36×10^{-09}	531	531	8.16×10^{-10}	0.02

APPENDIX D: THE FITNESS-INDUCED TWO-STAR MODEL AND ITS NUMERICAL RESOLUTION

In order to solve the system of equations defining the fit2SM, an appropriate vector of initial conditions needs to be chosen. In order to solve the dcGM, we have chosen $z_0 = 1$; in order to solve the fit2SM, we have chosen $\kappa_i^{(0)} = \kappa_i^{\text{dcGM}}$, $z_0 = z_{\text{dcGM}}$, and $y_0 = 1$. As a stopping criterium, we have adopted a condition on the infinite norm of the vector of differences between the values of the parameters at subsequent iterations, i.e., $\max\{|\Delta z|, |\Delta y|\} < 10^{-12}$. The accuracy of our method in estimating the constraints has been evaluated by computing the *maximum relative errors* defined as $\text{MRE}_L = |L^* - \langle L \rangle|/L^*$ and $\text{MRE}_S = |S^* - \langle S \rangle|/S^*$. Table I shows the time employed by our algorithm to converge as well as its accuracy in reproducing the constraints defining it. Overall, our method is fast and accurate: the numerical errors never exceed $O(10^{-1})$ and the time employed to achieve such an accuracy is always less than a minute.

A natural question arises, i.e., does employing the “single-iteration” solution lead to significantly different results? To answer this question, we have compared the “single-iteration” (si) solutions with the self-consistent (sc) ones by calculating the maximum relative error concerning the average BIC values for each timescale: one finds that $|\text{BIC}_{\text{sc}} - \text{BIC}_{\text{si}}|/\text{BIC}_{\text{sc}} = 0.0016$ at the daily timescale, $|\text{BIC}_{\text{sc}} - \text{BIC}_{\text{si}}|/\text{BIC}_{\text{sc}} = 0.0011$ at the weekly timescale, $|\text{BIC}_{\text{sc}} - \text{BIC}_{\text{si}}|/\text{BIC}_{\text{sc}} = 0.0003$ at the monthly timescale, $|\text{BIC}_{\text{sc}} - \text{BIC}_{\text{si}}|/\text{BIC}_{\text{sc}} = 0.0008$ at the quarterly timescale, and $|\text{BIC}_{\text{sc}} - \text{BIC}_{\text{si}}|/\text{BIC}_{\text{sc}} = 0.0008$ at the yearly timescale (see also Fig. 7).

Let us now discuss an alternative way of determining the parameters of the fit2SM. First, let us write the log-likelihood as

$$\mathcal{L}_{\text{fit2SM}} = \sum_i \sum_{j(>i)} [a_{ij} \ln(zs_i s_j y^{\kappa_i + \kappa_j}) - \ln(1 + zs_i s_j y^{\kappa_i + \kappa_j})] \quad (\text{D1})$$

and, then, maximize it with respect to z and y . Upon doing so, we derive the two nonlinear coupled equations reading

$$L^* = \sum_i \sum_{j(>i)} p_{ij}^{\text{fit2SM}} = \langle L \rangle \quad (\text{D2})$$

and

$$\tau^* = \sum_i \sum_{j(>i)} a_{ij}(\kappa_i + \kappa_j) = \sum_i \sum_{j(>i)} p_{ij}^{\text{fit2SM}}(\kappa_i + \kappa_j) = \langle \tau \rangle; \quad (\text{D3})$$

given that

$$S = \frac{1}{2} \sum_i \sum_{j(>i)} a_{ij}(k_i + k_j) - L = \frac{T}{2} - L, \quad (\text{D4})$$

reproducing L and (a proxy of) T amounts at reproducing L and (a proxy of) S . Since, however, we are interested in reproducing the empirical value of S , we have adopted the so-called *method of moments*, imposing $\langle S \rangle = S^*$ in a (more) direct fashion.

In case the total number of two-stars were not directly accessible, one could exploit the relationship between L and S , replacing S^* with the value $S = aL^b$. For what concerns eMID, the fitted values read $a_{\text{daily}} \simeq 0.36$, $b_{\text{daily}} \simeq 1.59$; $a_{\text{weekly}} \simeq 0.36$, $b_{\text{weekly}} \simeq 1.61$; $a_{\text{monthly}} \simeq 0.47$, $b_{\text{monthly}} \simeq 1.59$; $a_{\text{quarterly}} \simeq 0.67$, $b_{\text{quarterly}} = 1.56$; and $a_{\text{yearly}} \simeq 0.69$, $b_{\text{yearly}} = 1.56$.

APPENDIX E: SPECTRAL RADIUS

We now assess the ability of the considered models to reproduce the empirical value of a network spectral radius: To this aim, we pose ourselves within the controlled framework described in the main text.

Figure 8 shows the distribution of the spectral radius induced by the $M = 10^3$ configurations sampled from the Chung-Lu model calibrated on the eMID snapshot corresponding to the day 2010-07-19, the average $\langle \lambda_1 \rangle \simeq 9.36$ being remarkably close to $\pi_1 \simeq \overline{k^2}/\overline{k} = 9.35$. Besides, we

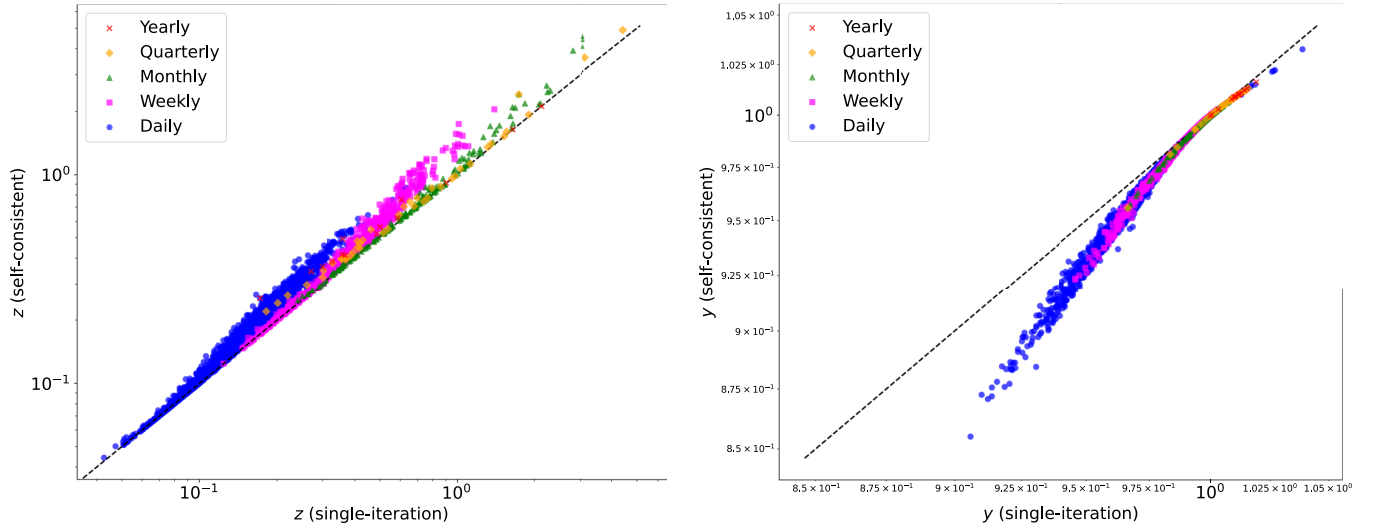


FIG. 7. Comparison between the “single-iteration” solutions and the self-consistent ones for all the years of our dataset, at each timescale.

scatter the estimation of λ_1 obtained from each of the $M = 10^3$, sampled configurations versus the corresponding empirical value: while the dcGM generally underestimates the latter,

the UBCM overestimates it, as a consequence of its tendency to overestimate the variance of the degree distribution; the fit2SM, instead, displays the most accurate results.

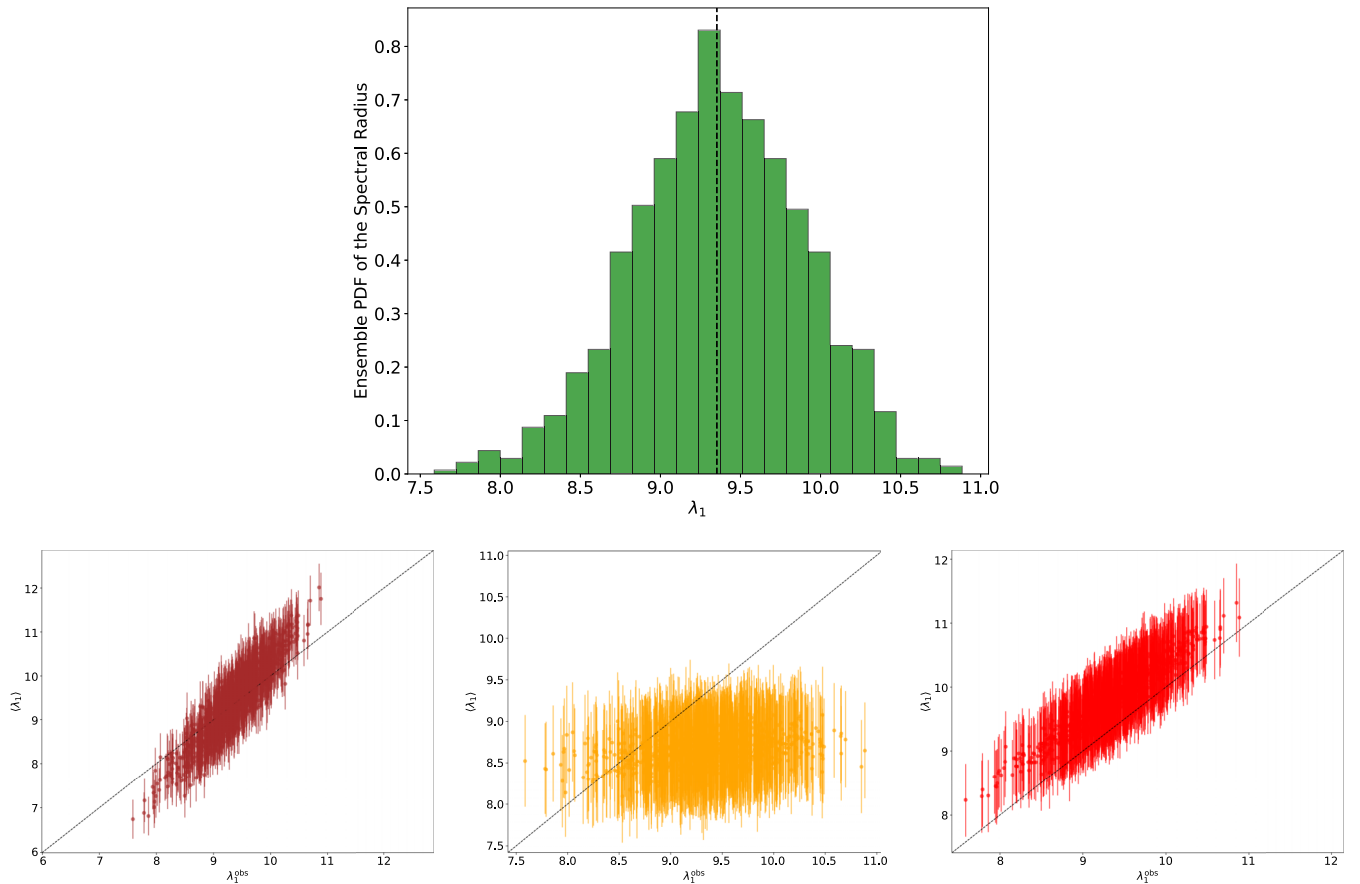


FIG. 8. Top panel: distribution of the spectral radius induced by $M = 10^3$ configurations sampled from the Chung-Lu model defined by the topology of the eMID snapshot corresponding to the day 2010-07-19. Bottom panels: estimations of λ_1 , obtained from each of the M , sampled configurations, scattered vs the corresponding empirical values, for the UBCM (red), the dcGM (yellow), and the fit2SM (brown). Vertical bars indicate the standard deviation accompanying the estimation carried out on the specific configuration. While the dcGM generally underestimates the spectral radius of the generative model and the UBCM overestimates it, the fit2SM captures it quite accurately.

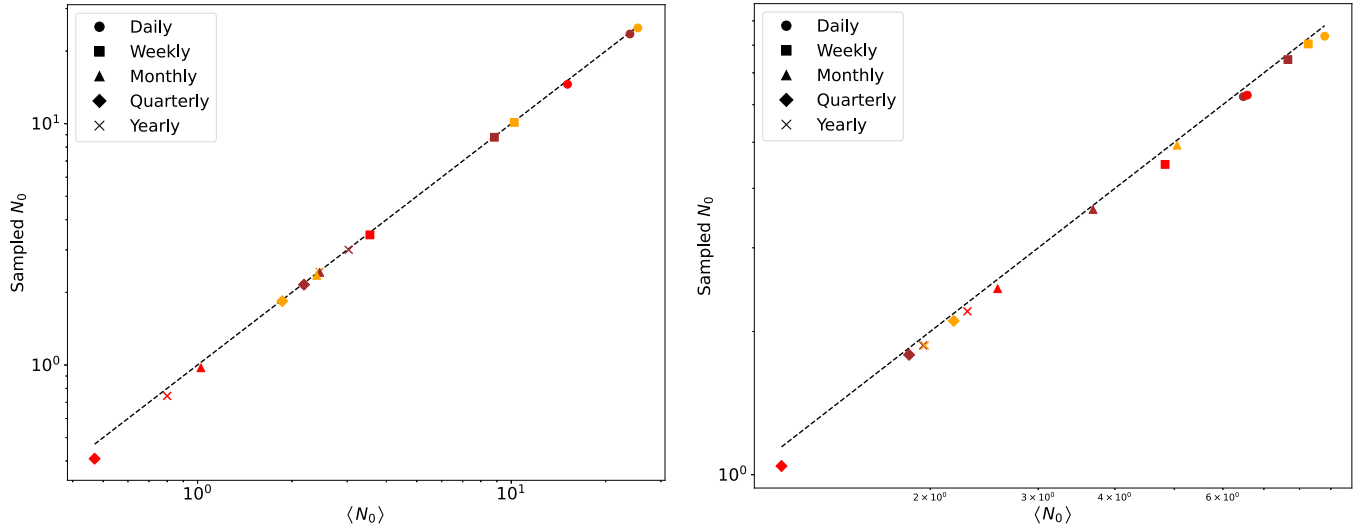


FIG. 9. Comparison between the ensemble averages of the total number of isolated nodes with their analytical counterparts for the weekly, monthly, quarterly, and yearly aggregation levels in 1999 (left) and 2012 (right): overall, the estimations provided by the UBCM (red), the dcGM (yellow), and the fit2SM (brown) are very accurate for each timescale. The quarter, month, week, and day shown for each year correspond to those reported in Table I.

APPENDIX F: PROBABILITY OF OBSERVING AT LEAST ONE ISOLATED NODE AND EXPECTED NUMBER OF ISOLATED NODES

Let us consider a sparse network. As such, dyads are independent and node i is isolated with probability

$$q_i^0 = P(k_i = 0) = \prod_{j(\neq i)} (1 - p_{ij}) \simeq \prod_{j(\neq i)} e^{-p_{ij}} = e^{-\langle k_i \rangle}; \quad (\text{F1})$$

as a consequence, no node is isolated with probability

$$q^* = \prod_i (1 - q_i^0) = \prod_i \left[1 - \prod_{j(\neq i)} (1 - p_{ij}) \right] \simeq \prod_i \left[1 - \prod_{j(\neq i)} e^{-p_{ij}} \right] = \prod_i [1 - e^{-\langle k_i \rangle}] \quad (\text{F2})$$

and at least one node is isolated with probability

$$q^0 = 1 - q^* = 1 - \prod_i (1 - q_i^0) = 1 - \prod_i \left[1 - \prod_{j(\neq i)} (1 - p_{ij}) \right] \simeq 1 - \prod_i \left[1 - \prod_{j(\neq i)} e^{-p_{ij}} \right] = 1 - \prod_i [1 - e^{-\langle k_i \rangle}]. \quad (\text{F3})$$

Let us now compare

$$q_{\text{UBCM}}^0 = 1 - \prod_i [1 - e^{-\langle k_i \rangle}] \quad \text{and} \quad \langle N \rangle_{\text{UBCM}}^0 = \sum_i q_i^0 = \sum_i e^{-\langle k_i \rangle}_{\text{UBCM}} \quad (\text{F4})$$

with

$$q_{\text{dcGM}}^0 = 1 - \prod_i [1 - e^{-\langle k_i \rangle_{\text{dcGM}}}] \quad \text{and} \quad \langle N \rangle_{\text{dcGM}}^0 = \sum_i q_i^0 = \sum_i e^{-\langle k_i \rangle_{\text{dcGM}}} \quad (\text{F5})$$

and with

$$q_{\text{fit2SM}}^0 = 1 - \prod_i [1 - e^{-\langle k_i \rangle_{\text{fit2SM}}}] \quad \text{and} \quad \langle N \rangle_{\text{fit2SM}}^0 = \sum_i q_i^0 = \sum_i e^{-\langle k_i \rangle_{\text{fit2SM}}} \quad (\text{F6})$$

across our temporal snapshots. Figure 9 compares the ensemble averages of the total number of isolated nodes with their analytical counterparts for the weekly, monthly, quar-

terly, and yearly aggregation levels in 1999 and 2012: Overall, the estimations provided above are very accurate at each timescale. Figure 10, instead, compares the performances of

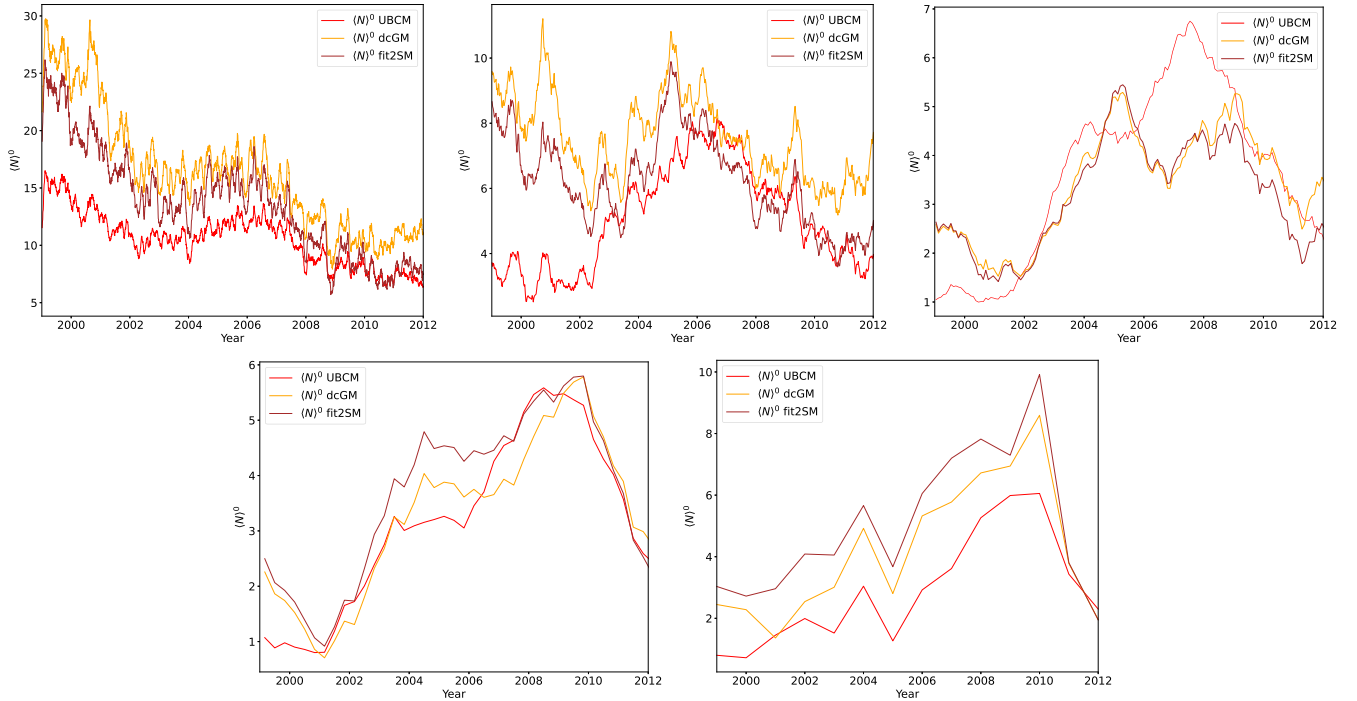


FIG. 10. Analysis of eMID at (from top-left to bottom-right) the daily, weekly, monthly, quarterly, and yearly timescales. According to the expected number of isolated nodes, the fit2SM outperforms the dcGM at the daily and weekly timescales and compete with the dcGM and the UBCM at the monthly and quarterly timescales. Trends have been smoothed via a rolling average over the points $[t - 10, t + 10]$ for the daily aggregation, $[t - 7, t + 7]$ for the weekly aggregation, $[t - 5, t + 5]$ for the monthly aggregation, $[t - 3, t + 3]$ for the quarterly aggregation.

the UBCM, the dcGM, and the fit2SM in “producing” configurations being characterized by a certain number of isolated nodes: as we have already observed, the sparser the configuration, the better the performance of the fit2SM—in these

cases, the fit2SM outperforms the dcGM and compete with the UBCM on at least a portion of the dataset, by allowing a smaller number of isolated nodes to appear on sampled configurations.

- [1] V. Colizza, A. Barrat, M. Barthélemy, and A. Vespignani, The role of the airline transportation network in the prediction and predictability of global epidemics, *Proc. Natl. Acad. Sci. USA* **103**, 2015 (2006).
- [2] A. Barrat, M. Barthélemy, and A. Vespignani, *Dynamical Processes on Complex Networks* (Cambridge University Press, Cambridge, UK, 2008).
- [3] M. Newman, *Networks: An Introduction* (Oxford University Press, Oxford, UK, 2010).
- [4] R. Pastor-Satorras, C. Castellano, P. Van Mieghem, and A. Vespignani, Epidemic processes in complex networks, *Rev. Mod. Phys.* **87**, 925 (2015).
- [5] T. Squartini, I. van Lelyveld, and D. Garlaschelli, Early-warning signals of topological collapse in interbank networks, *Sci. Rep.* **3**, 3357 (2013).
- [6] S. Battiston, *et al.*, Complexity theory and financial regulation, *Science* **351**, 818 (2016).
- [7] M. Bardoscia, S. Battiston, F. Caccioli, and G. Caldarelli, Pathways towards instability in financial networks, *Nat. Commun.* **8**, 14416 (2017).
- [8] V. Macchiati, E. Marchese, P. Mazzarisi, D. Garlaschelli, and T. Squartini, Spectral signatures of structural change in financial networks, *Chaos, Solitons Fractals* **193**, 116065 (2025).
- [9] E. T. Jaynes, Information theory and statistical mechanics, *Phys. Rev.* **106**, 620 (1957).
- [10] J. Park and M. E. J. Newman, Statistical mechanics of networks, *Phys. Rev. E* **70**, 066117 (2004).
- [11] D. Garlaschelli and M. I. Loffredo, Maximum likelihood: Extracting unbiased information from complex networks, *Phys. Rev. E* **78**, 015101 (2008).
- [12] T. Squartini and D. Garlaschelli, Analytical maximum-likelihood method to detect patterns in real networks, *New J. Phys.* **13**, 083001 (2011).
- [13] F. Saracco, R. Di Clemente, A. Gabrielli, and T. Squartini, Randomizing bipartite networks: The case of the World Trade Web, *Sci. Rep.* **5**, 10595 (2015).
- [14] T. Squartini and D. Garlaschelli, *Maximum-Entropy Networks* (Springer, Berlin, 2017).
- [15] G. Cimini, *et al.*, The statistical physics of real-world networks, *Nat. Rev. Phys.* **1**, 58 (2019).
- [16] A. Fronczak, P. Fronczak, and J. A. Holyst, Statistical mechanics of the international trade network, *Phys. Rev. E* **73**, 016108 (2006).
- [17] G. Bianconi, Entropy of network ensembles, *Phys. Rev. E* **79**, 036114 (2009).

- [18] A. Fronczak and P. Fronczak, Statistical mechanics of the international trade network: Structural correlations and modeling, *Phys. Rev. E* **85**, 056113 (2012).
- [19] T. Squartini, R. Mastrandrea, and D. Garlaschelli, Unbiased sampling of network ensembles, *New J. Phys.* **17**, 023052 (2015).
- [20] M. Molloy and B. Reed, A critical point for random graphs with a given degree sequence, *Random Structures & Algorithms* **6**, 161 (1995).
- [21] Y. Artzy-Randrup and L. Stone, Generating uniformly distributed random networks, *Physica A* **364**, 513 (2006).
- [22] C. I. Del Genio, H. Kim, Z. Toroczkai, and K. E. Bassler, Efficient and exact sampling of simple graphs with given arbitrary degree sequence, *PLoS One* **5**, e10012 (2010).
- [23] J. Blitzstein and P. Diaconis, A sequential importance sampling algorithm for generating random graphs with prescribed degrees, *Internet Math.* **6**, 489 (2011).
- [24] J. Kim, Z. Toroczkai, P. L. Erdős, I. Miklós, and I. Kondor, Network-based prediction of unobserved edges in complex networks, *New J. Phys.* **14**, 023012 (2012).
- [25] E. S. Roberts and A. C. C. Coolen, Unbiased degree-preserving randomization of directed binary networks, *Phys. Rev. E* **85**, 046103 (2012).
- [26] T. Squartini, G. Caldarelli, G. Cimini, A. Gabrielli, and D. Garlaschelli, Reconstruction methods for networks: The case of economic and financial systems, *Phys. Rep.* **757**, 1 (2018).
- [27] G. Cimini, T. Squartini, D. Garlaschelli, and A. Gabrielli, Systemic risk analysis on reconstructed economic and financial networks, *Sci. Rep.* **5**, 15758 (2015).
- [28] A. Almog, R. Bird, and D. Garlaschelli, Enhanced gravity model of trade: Reconciling macroeconomic and network models, *Front. Phys.* **7**, 55 (2019).
- [29] S. Hooijmaaijers and G. Buiten, Estimating the Dutch Interfirm Trade Network with Commodity Breakdown, Statistics Netherlands, Technical Report (2019).
- [30] C. E. S. Mattsson, *et al.*, Functional structure in production networks, *Front. Big Data* **4**, 666712 (2021).
- [31] L. N. Ialongo, *et al.*, Reconstructing firm-level interactions in the Dutch input–output network, *Sci. Rep.* **12**, 11847 (2022).
- [32] G. Cimini, R. Mastrandrea, and T. Squartini, *Reconstructing Networks, Elements in the Structure and Dynamics of Complex Networks* (Cambridge University Press, Cambridge, 2021).
- [33] S. Battiston, M. Puliga, R. Kaushik, P. Tasca, and G. Caldarelli, Debrank: Too central to fail? Financial networks, the Fed and systemic risk, *Sci. Rep.* **2**, 541 (2012).
- [34] R. Pastor-Satorras and A. Vespignani, Epidemic spreading in scale-free networks, *Phys. Rev. Lett.* **86**, 3200 (2001).
- [35] V. Sood and S. Redner, Voter model on heterogeneous graphs, *Phys. Rev. Lett.* **94**, 178701 (2005).
- [36] V. Sood, T. Antal, and S. Redner, Voter models on heterogeneous networks, *Phys. Rev. E* **77**, 041121 (2008).
- [37] G. Iori, G. D. Masi, O. Precup, G. Gabbi, and G. Caldarelli, A network analysis of the Italian overnight money market, *J. Econ. Dyn. Control* **32**, 259 (2008).
- [38] K. Finger, D. Fricke, and T. Lux, *Network Analysis of the E-MID Overnight Money Market: The Informational Value of Different Aggregation Levels for Intrinsic Dynamic Processes*, Kiel Working Paper No. 1782 (Kiel Institute for the World Economy, 2012).
- [39] J. Park and M. E. J. Newman, Solution of the two-star model of a network, *Phys. Rev. E* **70**, 066146 (2004).
- [40] The asterisk indicates (the value of the quantities measured on) the observed configuration $\mathbf{A}^*$.
- [41] Limiting ourselves to the position $p_{ij} = xy^{k_i+k_j}/(1+xy^{k_i+k_j})$ would, in fact, lead to an inconsistency since the information on the degrees is not supposed to be available.
- [42] To ease numerical manipulations, we have rescaled the strengths, dividing them by their arithmetic mean.
- [43] N. Vallarano, M. J. Straka, G. Cimini, and T. Squartini, Fast and scalable likelihood maximization for exponential random graph models via GPU parallelization, *Sci. Rep.* **11**, 5725 (2021).
- [44] Recall that the quantile q of a set of (ordered) values indicates the percentage q of the smallest values. Given the empirical degree sequence and the model-based, expected one, we, thus, compute their quantile functions by sorting both in increasing order and scattering the r th order statistics $\langle k \rangle_r$ versus k_r , $r = 1, \dots, N$.
- [45] F. Chung and L. Lu, Connected components in random graphs with given expected degree sequences, *Ann. Comb.* **6**, 125 (2002).
- [46] R. A. Horn and C. R. Johnson, *Matrix Analysis* (Cambridge University Press, Cambridge, 1985).
- [47] <https://github.com/mattiamarzi/fit2SM>.