

Frontiers of Network Science

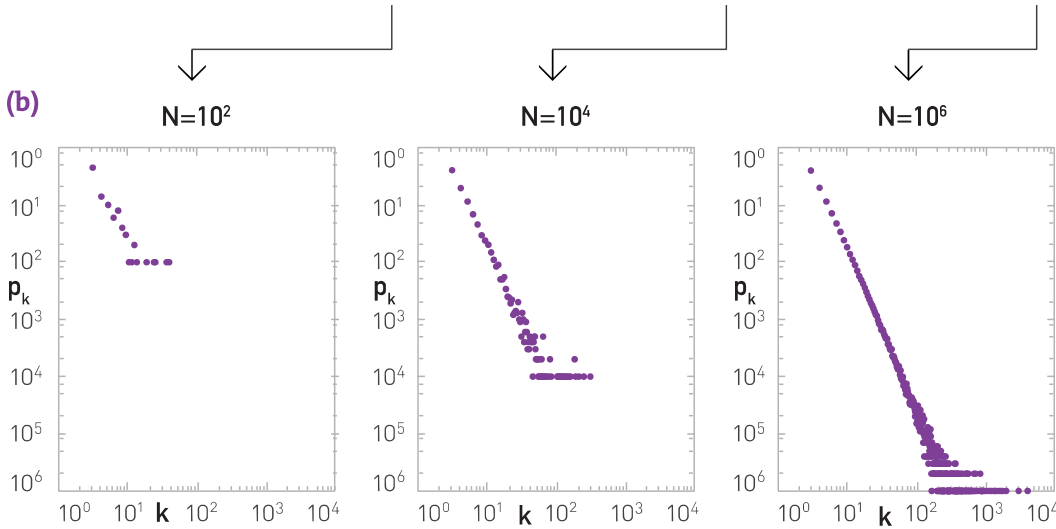
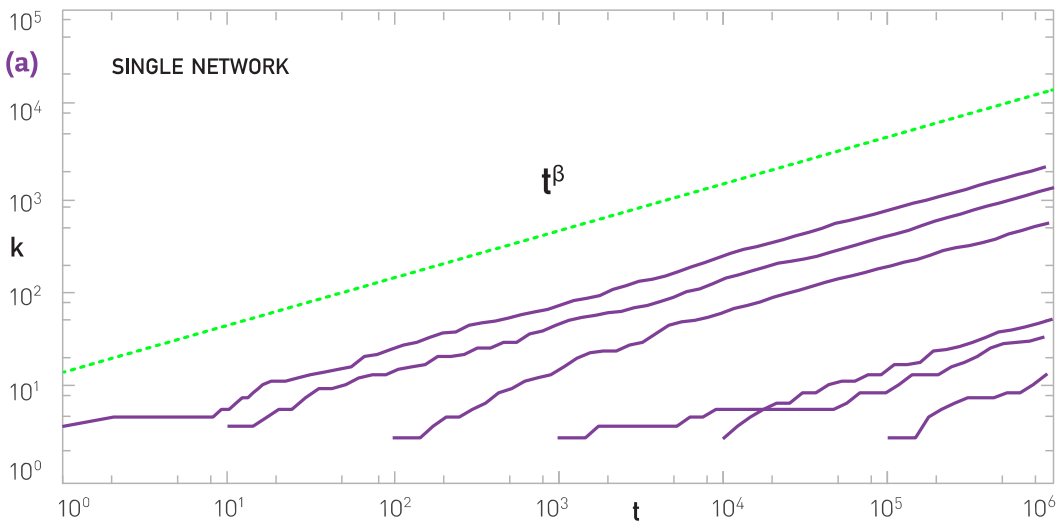
Fall 2021

Class 16: Degree Correlations **(Chapter 7 in Textbook)**

Boleslaw Szymanski

based on slides by
Albert-László Barabási
and Roberta Sinatra

Section 5.3



- The degree of each node increases following a power-law with the same dynamical exponent $\beta = 1/2$ (Figure 5.6a). Hence all nodes follow the same dynamical law.
- The growth in the degrees is sublinear (i.e. $\beta < 1$). This is a consequence of the growing nature of the Barabási-Albert model: Each new node has more nodes to link to than the previous node. Hence, with time the existing nodes compete for links with an increasing pool of other nodes.
- The earlier node i was added, the higher is its degree $k_i(t)$. Hence, hubs are large because they arrived earlier, a phenomenon called *first-mover advantage* in marketing and business.
- The rate at which the node i acquires new links is given by the derivative of (5.7)

$$\frac{dk_i(t)}{dt} = \frac{m}{2} \frac{1}{\sqrt{t_i t}} \tag{5.8}$$

indicating that in each time frame older nodes acquire more links (as they have smaller t_i). Furthermore the rate at which a node acquires links decreases with time as $t^{-1/2}$. Hence, fewer and fewer links go to a node.

Absence of growth and preferential attachment

MODEL A

growth

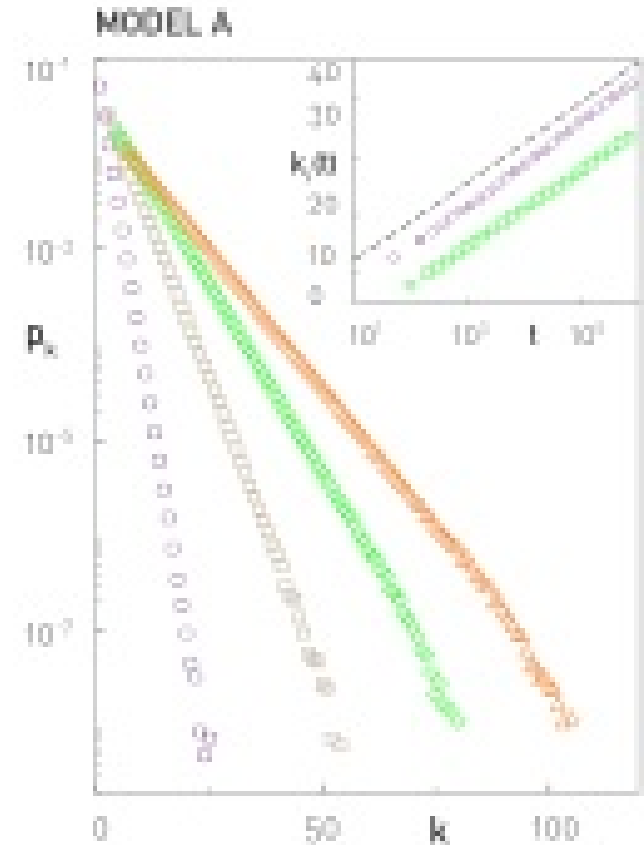
~~preferential attachment~~

$\Pi(k_i)$: uniform

$$\frac{\partial k_i}{\partial t} = A\Pi(k_i) = \frac{m}{m_0 + t - 1}$$

$$k_i(t) = m \ln\left(\frac{m_0 + t - 1}{m + t_i - 1}\right) + m$$

$$P(k) = \frac{e}{m} \exp\left(-\frac{k}{m}\right) \sim e^{-k}$$



MODEL B

~~growth~~

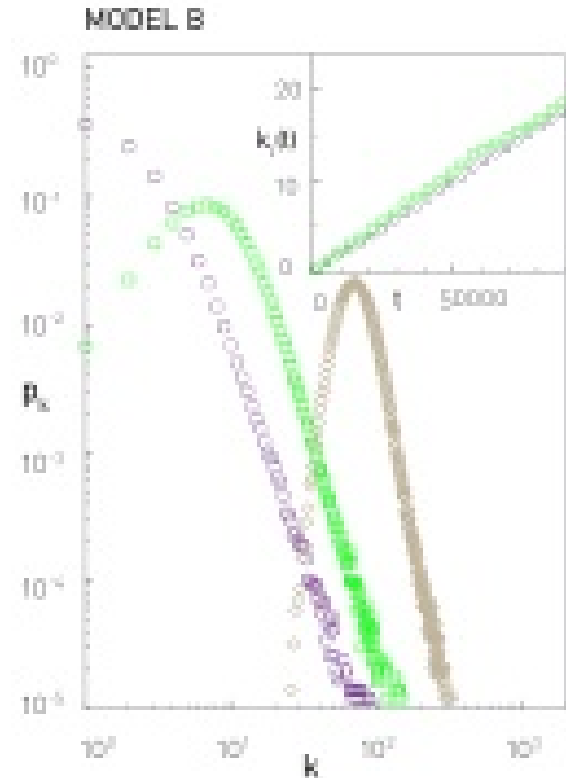
preferential attachment

$$\frac{\partial k_i}{\partial t} = A\Pi(k_i) + \frac{1}{N} = \frac{N}{N-1} \frac{k_i}{2t} + \frac{1}{N}$$

$$k_i(t) = \frac{2(N-1)}{N(N-2)}t + Ct^{\frac{N}{2(N-1)}} \sim \frac{2}{N}t$$

p_k : power law (initially) $\rightarrow$

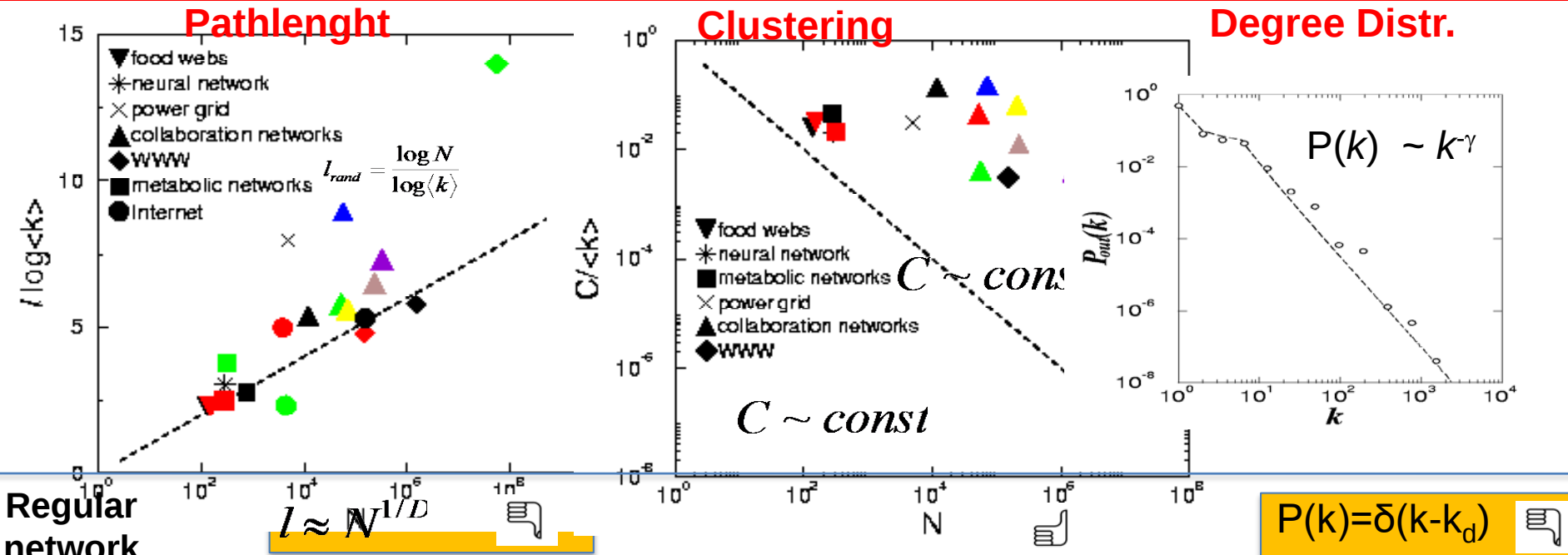
$\rightarrow$ Gaussian $\rightarrow$ Fully Connected



Do we need both growth and
preferential attachment?

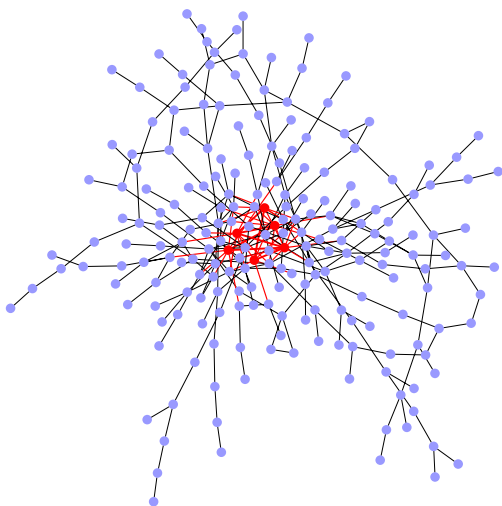
YEP

EMPIRICAL DATA FOR REAL NETWORKS



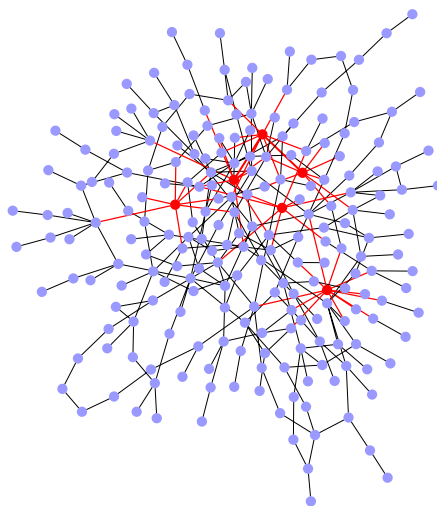
Regular network	$l \approx N^{1/D}$	$C \sim const$	$P(k) = \delta(k - k_d)$
Erdos-Renyi	$l_{rand} \approx \frac{\log N}{\log \langle k \rangle}$	$C_{rand} = p = \frac{\langle k \rangle}{N}$	$P(k) = e^{-\langle k \rangle} \frac{\langle k \rangle^k}{k!}$
Watts-Strogatz	$l_{rand} \approx \frac{\log N}{\log \langle k \rangle}$	$C \sim const$	Exponential
Barabasi-Albert			$P(k) \sim k^{-\gamma}$

DEGREE CORRELATIONS IN NETWORKS



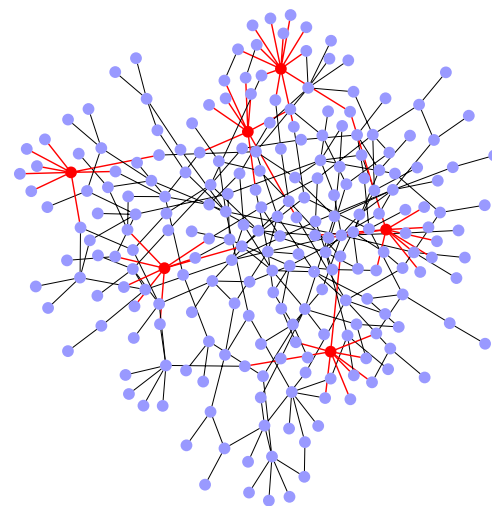
Assortative:

hubs show a tendency to link to each other.



Neutral:

nodes connect to each other with the expected random probabilities.



Disassortative:

Hubs tend to avoid linking to each other.

Quantifying degree correlations (three approaches):

- full statistical description (Maslov and Sneppen, Science 2001)
- degree correlation function (Pastor Satorras and Vespignani, PRL 2001)
- correlation coefficient (Newman, PRL 2002)

STATISTICAL DESCRIPTION

e_{jk} : probability to find a node with degree j and degree k at the two ends of a randomly selected edge

$$\sum_{j,k} e_{jk} = 1 \quad \sum_j e_{jk} = q_k$$

q_k : the probability to have a degree k node at the end of a link.

Where: $q_k = \frac{k p_k}{\langle k \rangle}$

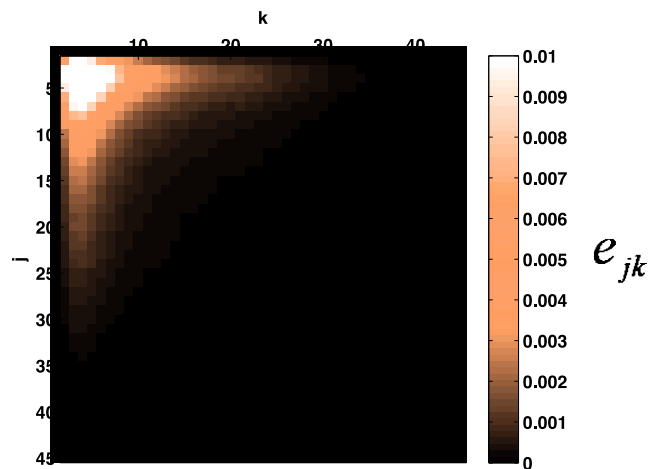
Probability to find a node at the end of a link is biased towards the more connected nodes, i.e. $q_k = C k p_k$, where C is a normalization constant. After normalization we find $C = 1/\langle k \rangle$, or $q_k = k p_k / \langle k \rangle$

If the network has no degree correlations:

$$e_{jk} = q_j q_k$$

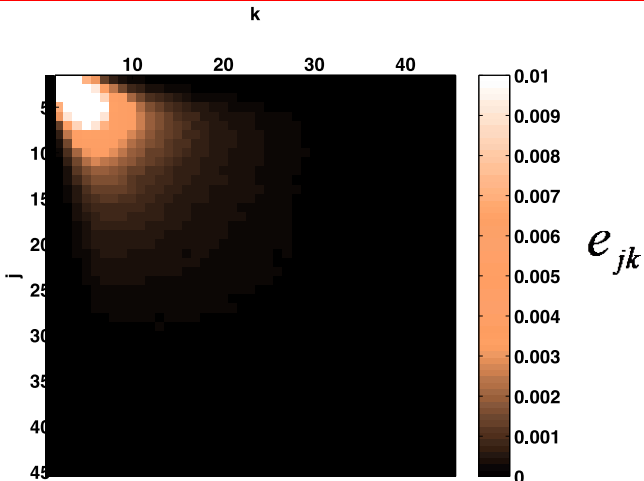
Deviations from this prediction are a signature of *degree correlation*.

EXAMPLE: e_{jk} FOR A SCALE-FREE NETWORK

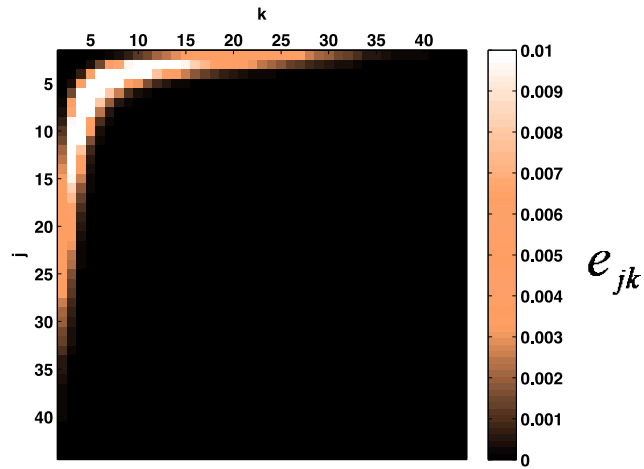


Neutral

Assortative:
More strength in the diagonal, hubs tend to link to each other.



Disassortative:
Hubs tend to connect to small nodes.



Each matrix is the average of a 100 independent scale-free networks, generated using the static model with $N=10^4$, $\gamma=2.5$ and $\langle k \rangle=3$.

EXAMPLE: e_{jk} FOR A SCALE-FREE NETWORK

*Perfectly assortative
network:*

$$e_{jk} = q_k \delta_{jk}$$

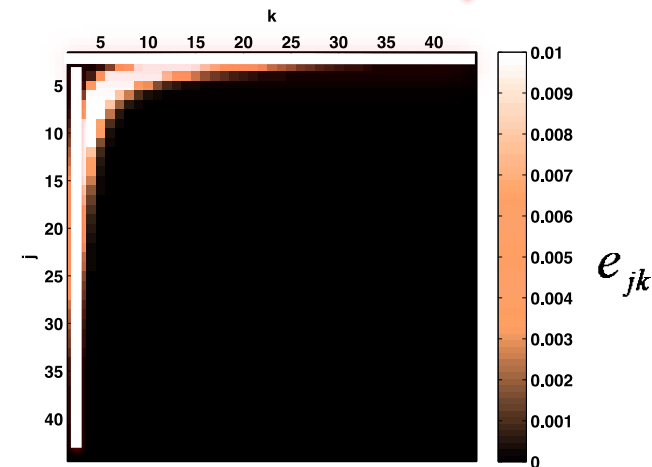
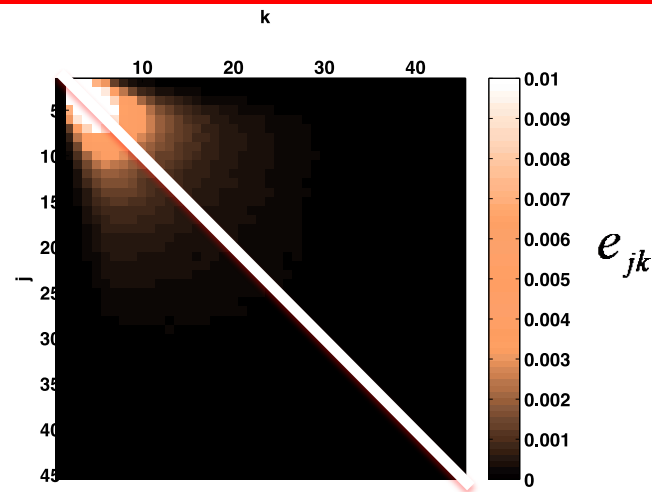
Assortative:

More strength in
the diagonal,
hubs tend to link
to each other.

*Perfectly
disassortative
network:*

Disassortative:

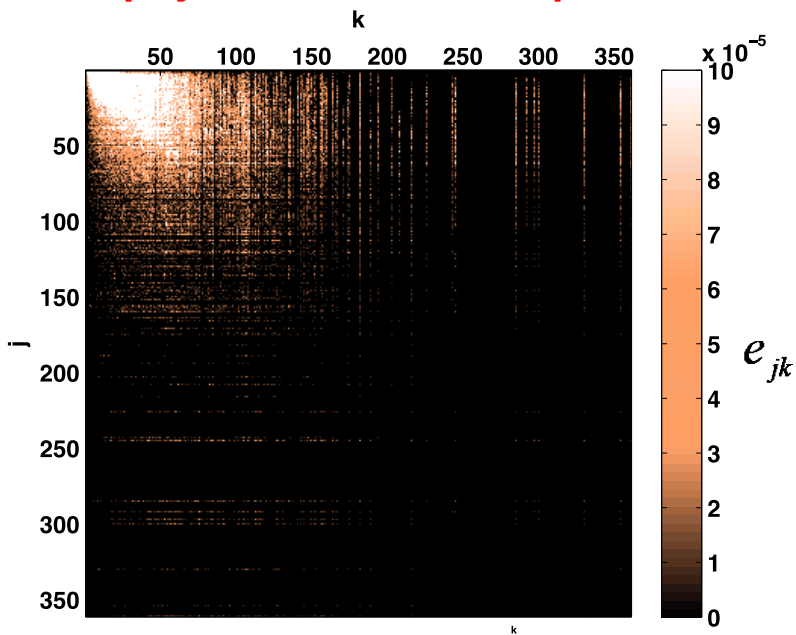
Hubs tend to
connect to small
nodes.



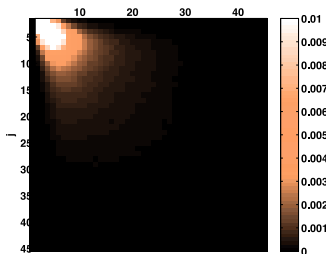
Each matrix is the average of a 100 independent scale-free networks, generated using the static model with $N=10^4$, $\gamma=2.5$ and $\langle k \rangle=3$.

REAL-WORLD EXAMPLES

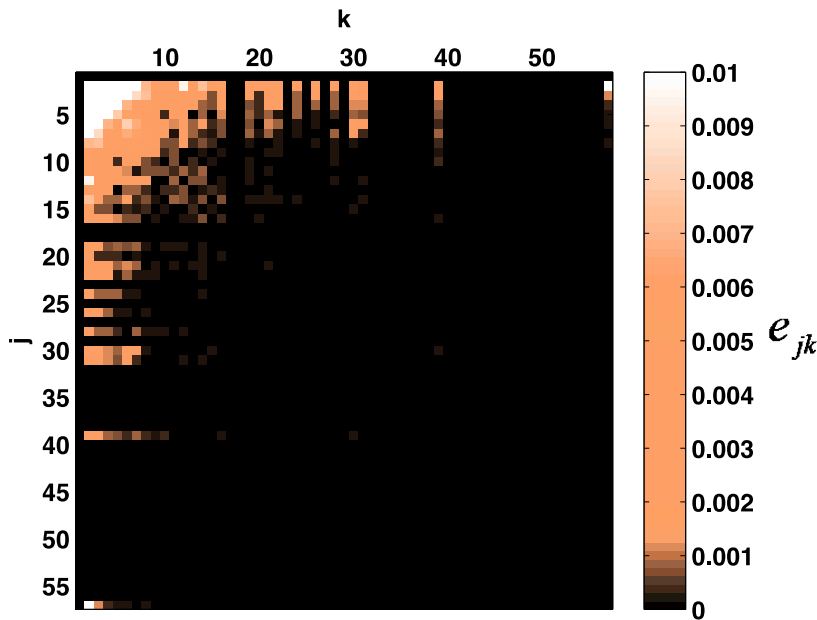
Astrophysics co-authorship network



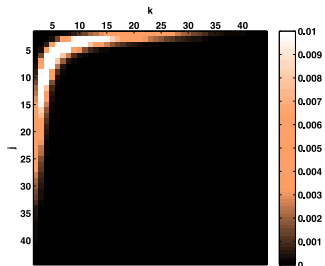
Assortative:
More strength in the diagonal, hubs tend to link to each other.



Yeast PPI

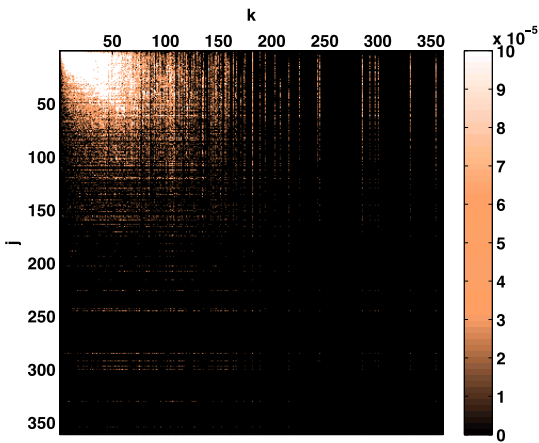


Disassortative:
Hubs tend to connect to small nodes.



PROBLEM WITH THE FULL STATISTICAL DESCRIPTION

(1) Difficult to extract information from a visual inspection of a matrix.



(2) Based on e_{jk} and hence requires a large number of elements to inspect:

$$\frac{k_{\max}(k_{\max}-1)}{2} - 1 - k_{\max}$$

Nr. of independent elements

Undirected network:
 $k_{\max} \times k_{\max}$ matrix

$\sum_{j,k} e_{jk} = 1$

$\sum_{j=1, k_{\max}} e_{jk} = q_k$

Constraints

We need to find a way to reduce the information contained in e_{jk}