



Lecture 6 - Graph partition and Node Centrality

Graph-Based Data Science For Networked Systems
ECE 5260 – ORIE 5735

Anna Scaglione

Cornell Tech

February 9, 2026



Content

- 1 Introduction
- 2 Betti number and Hodge Laplacian
- 3 Graph partitioning
- 4 The ranking of nodes: centrality measures
- 5 Conclusions

Outline



Cornell University

1 Introduction

2 Betti number and Hodge Laplacian

3 Graph partitioning

4 The ranking of nodes: centrality measures

5 Conclusions



Summary

- We introduced basic graph theory concepts, defining weighted and unweighted graphs $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ and graph classes
- We studied these graphs properties using:
 - The adjacency matrix A :
 - the basic algebraic structure to highlight nodes connectivity
 - The incidence matrix B :
 - the basic structure to differentiate nodal quantities and model the impact of edge flows
 - The Laplacian matrix L :
 - a combination of the two. Its spectrum is tied to graph connectivity
 - The Hodge Laplacian which explores the simplicial structure of the graph (model higher-order interactions)
- We introduced algorithms to solve network flow problems and graph partitioning



Summary

- We introduced basic graph theory concepts, defining weighted and unweighted graphs $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ and graph classes
- We studied these graphs properties using:
 - The adjacency matrix A :
 - the basic algebraic structure to highlight nodes connectivity
 - The incidence matrix B :
 - the basic structure to differentiate nodal quantities and model the impact of edge flows
 - The Laplacian matrix L :
 - a combination of the two. Its spectrum is tied to graph connectivity
 - The Hodge Laplacian which explores the simplicial structure of the graph (model higher-order interactions)
 - We introduced algorithms to solve network flow problems and graph partitioning



Summary

- We introduced basic graph theory concepts, defining weighted and unweighted graphs $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ and graph classes
- We studied these graphs properties using:
 - The adjacency matrix A :
 - the basic algebraic structure to highlight nodes connectivity
 - The incidence matrix B :
 - the basic structure to differentiate nodal quantities and model the impact of edge flows
 - The Laplacian matrix L :
 - a combination of the two. Its spectrum is tied to graph connectivity
 - The Hodge Laplacian which explores the simplicial structure of the graph (model higher-order interactions)
 - We introduced algorithms to solve network flow problems and graph partitioning



Summary

- We introduced basic graph theory concepts, defining weighted and unweighted graphs $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ and graph classes
- We studied these graphs properties using:
 - The adjacency matrix A :
 - the basic algebraic structure to highlight nodes connectivity
 - The incidence matrix B :
 - the basic structure to differentiate nodal quantities and model the impact of edge flows
 - The Laplacian matrix L :
 - a combination of the two. Its spectrum is tied to graph connectivity
 - The Hodge Laplacian which explores the simplicial structure of the graph (model higher-order interactions)
 - We introduced algorithms to solve network flow problems and graph partitioning



Summary

- We introduced basic graph theory concepts, defining weighted and unweighted graphs $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ and graph classes
- We studied these graphs properties using:
 - The adjacency matrix A :
 - the basic algebraic structure to highlight nodes connectivity
 - The incidence matrix B :
 - the basic structure to differentiate nodal quantities and model the impact of edge flows
 - The Laplacian matrix L :
 - a combination of the two. Its spectrum is tied to graph connectivity
 - The Hodge Laplacian which explores the simplicial structure of the graph (model higher-order interactions)
- We introduced algorithms to solve network flow problems and graph partitioning



Introduction

- From these matrices we have derived through algebra a set of graph properties and explored models that apply to networked systems that use them.
- Last time in particular we talked about the Laplacian, some results in spectral graph theory and introduced the Hodge Laplacian and introduced the problem of clustering a network and the spectral clustering method.
- We also defined the Hodge-Laplacian



Introduction

- From these matrices we have derived through algebra a set of graph properties and explored models that apply to networked systems that use them.
- Last time in particular we talked about the Laplacian, some results in spectral graph theory and introduced the Hodge Laplacian and introduced the problem of clustering a network and the spectral clustering method.
- We also defined the Hodge-Laplacian

Overview



- We will briefly see a key spectral property of the Hodge Laplacian $\mathbf{L}_k$
- We want to better explain the algorithms for graph partitioning using the Laplacian matrix $\mathbf{L}$
- We will take a first look first at network structural features (measures and metrics) through the notion of centrality

Outline



Cornell University

1 Introduction

2 Betti number and Hodge Laplacian

3 Graph partitioning

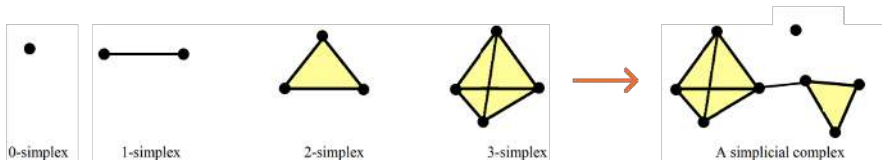
4 The ranking of nodes: centrality measures

5 Conclusions



Simplicial Complexes

- A **simplicial complex** x_K of order K is a collection of simplices S_k for $k = 0, \dots, K$.
- A k -simplex is a set of $k + 1$ vertices and, if the nodes are embedded in space, it corresponds to a k -dimensional polytope that is the convex hull of its $k + 1$ vertices. For example:
 - Nodes (0-simplices)
 - Edges (1-simplices)
 - Triangles (2-simplices)
 - Tetrahedrons (3-simplices)
- If $S_k \in x_K$, then all its faces S_{k-1} must also be in x_K .



Betti Numbers and the Hodge Laplacian in Graphs



Cornell University

- **Betti numbers** are topological invariants that describe the number of independent cycles in a space.
- In graph theory, they help understand the topology of a network.
- The **Hodge Laplacian** provides a spectral decomposition of graph structure and relates to Betti numbers via the topology of the graph.
- While the spectrum of $\mathbf{L}_0$, the regular Laplacian, is tied to the number of connected components of the graphs, for higher orders $\mathbf{L}_k$ provides information about the presence of *holes* and cycles in the network fabric.



Betti Numbers: Definition

Given a simplicial complex X , the n -th **Betti number** is:

$$b_n = \text{rank } H_n(X),$$

where $H_n(X)$ is the n -th homology group.

Intuition:

- b_0 : Number of connected components.
- b_1 : Number of independent cycles.
- b_2 : Number of independent voids / cavities (in higher dimensions)
-

Independent Cycles and the Cycle Space



Cycle space

Given an undirected graph $G = (V, E)$ with an arbitrary orientation of its edges, let $B_1 \in \mathbb{R}^{|V| \times |E|}$ be the node-edge incidence matrix. The *cycle space* of the graph is defined as

$$\mathcal{C} := \ker B_1.$$

where $\ker$ is the Kernel of the matrix B_1 . Vectors $f \in \mathcal{C}$, i.e. such that

$$B_1 f = 0,$$

represent edge flows that satisfy **flow conservation** at every node.

Independent Cycles and the Cycle Space



Cornell University

Independent cycles

A set of cycles is said to be *independent* if the corresponding vectors in $\ker \mathbf{B}_1$ are linearly independent. The number of independent cycles in the graph is therefore

$$b_1 = \dim(\ker \mathbf{B}_1),$$

which is the **first Betti** number of the graph.

Interpretation

Each independent cycle corresponds to one degree of freedom for circulating flow that does not accumulate at any node.

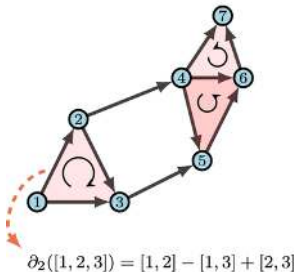
Equivalently, each edge added to a spanning tree creates exactly one new independent cycle.



Incidence Matrices

The relationships between simplices are captured by incidence matrices:

- B_1 : Node-to-edge incidence matrix.
- B_2 : Edge-to-triangle incidence matrix.
- ...
- B_k : Edge-to-k-simplexes incidence matrix.



$$B_1 = \begin{pmatrix} \begin{matrix} [1] \\ [2] \\ [3] \\ [4] \\ [5] \\ [6] \\ [7] \end{matrix} & \begin{matrix} [1] \\ [2] \\ [3] \\ [4] \\ [5] \\ [6] \\ [7] \end{matrix} & \begin{matrix} [2] \\ [3] \\ [4] \\ [5] \\ [6] \\ [7] \end{matrix} & \begin{matrix} [2] \\ [3] \\ [4] \\ [5] \\ [6] \\ [7] \end{matrix} & \begin{matrix} [3] \\ [4] \\ [5] \\ [6] \\ [7] \end{matrix} & \begin{matrix} [4] \\ [5] \\ [6] \\ [7] \end{matrix} & \begin{matrix} [4] \\ [5] \\ [6] \\ [7] \end{matrix} & \begin{matrix} [5] \\ [6] \\ [7] \end{matrix} & \begin{matrix} [6] \\ [7] \end{matrix} \end{pmatrix}$$

$$B_2 = \begin{pmatrix} \begin{matrix} [1, 2] \\ [1, 3] \\ [2, 3] \\ [2, 4] \\ [3, 5] \\ [4, 5] \\ [4, 6] \\ [4, 7] \\ [5, 6] \\ [6, 7] \end{matrix} & \begin{matrix} [1] \\ [2] \\ [3] \\ [4] \\ [5] \\ [6] \\ [7] \end{matrix} & \begin{matrix} [4] \\ [5] \\ [6] \\ [7] \end{matrix} \end{pmatrix}$$

They encode the boundary information between simplex orders.



From Cycles to Homology

Chains and boundary operators

For a simplicial complex X , let C_k denote the vector space of k -chains, i.e., formal linear combinations of k -simplices. The boundary operator

$$\partial_k : C_k \rightarrow C_{k-1}$$

maps a k -simplex to the signed sum of its $(k-1)$ -dimensional faces. In matrix form, ∂_k is represented by the incidence matrix B_k .

k -cycles and k -boundaries

- **k -cycles**: chains with no boundary

$$Z_k := \ker \partial_k = \ker B_k$$

- **k -boundaries**: chains that are boundaries of higher-order simplices

$$B_k := \operatorname{im} \partial_{k+1} = \operatorname{im} B_{k+1}$$



Hodge Laplacian

- The Hodge Laplacian at level k is defined as:

$$\mathbf{L}_k = \mathbf{B}_k^T \mathbf{B}_k + \mathbf{B}_{k+1} \mathbf{B}_{k+1}^T$$

the regular Laplacian is $\mathbf{L}_0$ as $\mathbf{B}_0 = \mathbf{0}$

- Decomposed into:

- Lower Laplacian: $\mathbf{L}_k^l = \mathbf{B}_k^T \mathbf{B}_k$

- Upper Laplacian: $\mathbf{L}_k^u = \mathbf{B}_{k+1} \mathbf{B}_{k+1}^T$

- The kernel of $\mathbf{L}_k$ gives Betti numbers (the number of zero eigenvalues of $\mathbf{L}_k$):

$$b_k = \dim(\ker \mathbf{L}_k) = \nu(\mathbf{L}_k). \quad (1)$$



Homology Groups and Betti Numbers

The k -th homology group

The k -th homology group is defined as

$$H_k := \mathbf{Z}_k / \mathbf{B}_k,$$

i.e., the space of k -cycles modulo those that are boundaries of $(k+1)$ -simplices.

Betti numbers

The k -th Betti number is the dimension of the homology group:

$$b_k := \dim H_k.$$

It counts the number of *independent k -dimensional holes* in the simplicial complex.

Betti Numbers and the Genus of a Surface



Cornell University

Two different but related notions

- **Betti numbers** are *intrinsic* topological invariants: they depend only on the connectivity of the simplicial complex.
- The **genus** g is an *extrinsic* geometric notion: it depends on how a graph is embedded on a surface.

For a graph viewed as a 1-dimensional simplicial complex:

- b_1 counts the total number of independent cycles.
- The genus g counts the minimum number of *handles* required to embed the graph without edge crossings.

Key message

A graph may have many independent cycles ($b_1 \gg 0$) and still be planar ($g = 0$). Genus measures how *entangled* the cycles are, not how many there are.

Euler Characteristic: Linking Cycles and Genus



Euler characteristic of a surface embedding

If a connected graph is embedded on a surface of genus g , then

$$\chi = |\mathcal{V}| - |\mathcal{E}| + |\mathcal{F}| = 2 - 2g,$$

where $|\mathcal{F}|$ is the number of faces induced by the embedding and χ is called Euler characteristic. Since for a connected graph, $b_1 = \ker \mathbf{B}_1 = |\mathcal{E}| - |\mathcal{V}| + 1$. Combining the two expressions yields:

$$b_1 = |\mathcal{F}| + 2g - 1.$$

Interpretation:

- b_1 counts *all* independent cycles.
- Faces account for cycles that can be “filled in”.
- Genus counts the cycles that *cannot* be removed without introducing crossings.



Betti Numbers in Graphs

For a graph G :

- $b_0 = \#$ of connected components.
- $b_1 = \#$ of independent cycles (from cycle space of the graph).

Computation from Hodge Laplacian:

- $b_0 = \dim(\ker \mathbf{L}_0)$ (0th Laplacian, connected components).
- $b_1 = \dim(\ker \mathbf{L}_1)$ (1st Laplacian, cycle space).



Outline

- 1 Introduction
- 2 Betti number and Hodge Laplacian
- 3 Graph partitioning**
- 4 The ranking of nodes: centrality measures
- 5 Conclusions

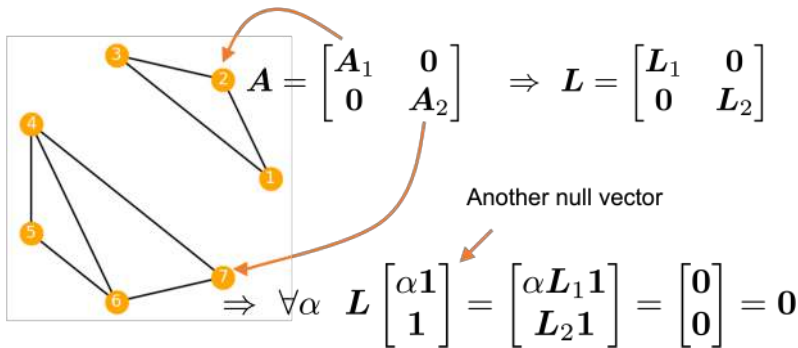


Recall what happens to disconnected graphs

The basic theorem of spectral graph theory is:

Theorem (M. Fiedler)

The **nullity** (number of zero eigenvalues) of L equals the number of connected components of a graph.

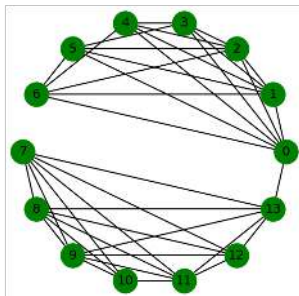
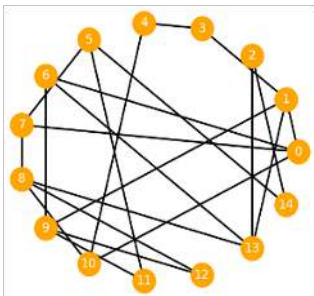




Graph partitioning and communities

- Even if the graph is connected, the second smallest eigenvalue of the Laplacian (called **algebraic connectivity** or **Fiedler eigenvalue**) $\lambda_{N-1} > 0$ can be very small in two cases:

- 1 When the network is sparsely connected throughout (it has low density $\delta = 2 \frac{|\mathcal{E}|}{|\mathcal{V}|(|\mathcal{V}|-1)}$), e.g. cycle graph when $|\mathcal{V}| \gg 1$
- 2 Networks that are dense, due to dense clusters, which are loosely connected among each other.

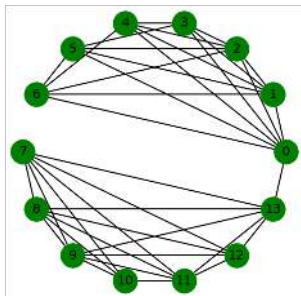
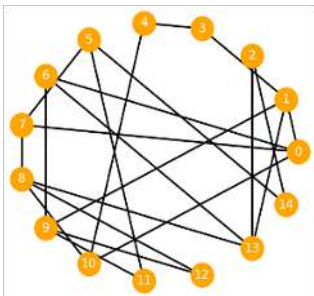




Graph partitioning and communities

- Even if the graph is connected, the second smallest eigenvalue of the Laplacian (called **algebraic connectivity** or **Fiedler eigenvalue**) $\lambda_{N-1} > 0$ can be very small in two cases:

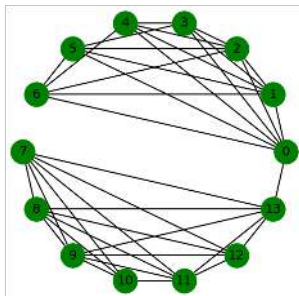
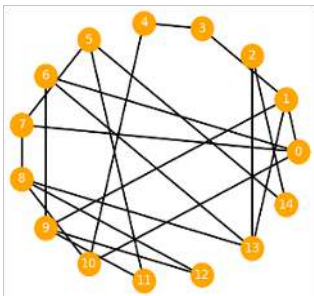
- 1 When the network is sparsely connected throughout (it has low density $\delta = 2 \frac{|\mathcal{E}|}{|\mathcal{V}|(|\mathcal{V}|-1)}$), e.g. cycle graph when $|\mathcal{V}| \gg 1$
- 2 Networks that are dense, due to dense clusters, which are loosely connected among each other.





Graph partitioning and communities

- Even if the graph is connected, the second smallest eigenvalue of the Laplacian (called **algebraic connectivity** or **Fiedler eigenvalue**) $\lambda_{N-1} > 0$ can be very small in two cases:
 - 1 When the network is sparsely connected throughout (it has low density $\delta = 2 \frac{|\mathcal{E}|}{|\mathcal{V}|(|\mathcal{V}|-1)}$), e.g. cycle graph when $|\mathcal{V}| \gg 1$
 - 2 Networks that are dense, due to dense clusters, which are loosely connected among each other.





Graph with dense sections (clusters)

- For instance, there could be two portions, corresponding to the blocks A_1 and A_2 , that are dense but not well connected:

$$A = \begin{bmatrix} A_1 & B_{12} \\ B_{21} & A_2 \end{bmatrix} \quad \|B_{12}\|_0 \ll |\mathcal{V}_1|, \quad \|B_{21}\|_0 \ll |\mathcal{V}_2|$$

- In fact, we can view the adjacency matrix of this graph as a *perturbed version of that of a disconnected graph* A^{disc} :

$$A = \begin{bmatrix} A_1 & 0 \\ 0 & A_2 \end{bmatrix} + \overbrace{\begin{bmatrix} 0 & B_{12} \\ B_{21} & 0 \end{bmatrix}}^{\delta A \text{ which is sparse}} \approx \overbrace{\begin{bmatrix} A_1 & 0 \\ 0 & A_2 \end{bmatrix}}^{A^{\text{disc}}}$$



Graph with dense sections (clusters)

- For instance, there could be two portions, corresponding to the blocks A_1 and A_2 , that are dense but not well connected:

$$A = \begin{bmatrix} A_1 & B_{12} \\ B_{21} & A_2 \end{bmatrix} \quad \|B_{12}\|_0 \ll |\mathcal{V}_1|, \quad \|B_{21}\|_0 \ll |\mathcal{V}_2|$$

- In fact, we can view the adjacency matrix of this graph as a *perturbed version of that of a disconnected graph* A^{disc} :

$$A = \begin{bmatrix} A_1 & 0 \\ 0 & A_2 \end{bmatrix} + \overbrace{\begin{bmatrix} 0 & B_{12} \\ B_{21} & 0 \end{bmatrix}}^{\delta A \text{ which is sparse}} \approx \overbrace{\begin{bmatrix} A_1 & 0 \\ 0 & A_2 \end{bmatrix}}^{A^{\text{disc}}}$$



Null vectors of the Laplacian

- A similar partition can be done for the Laplacian:

$$\mathbf{L} = \mathbf{L}^{\text{disc}} + \delta \mathbf{L}$$

where $\delta \mathbf{L} = -\delta \mathbf{A}$, since $\mathbf{L} = \text{diag}(\mathbf{d}) - \mathbf{A}$, so the off diagonal elements of $\mathbf{L}$ have simply a $-$ sign compared to those of $\mathbf{A}$.

- We know that $\mathbf{L}^{\text{disc}}$ has Fiedler eigenvalue $\lambda_{N-1}^{\text{disc}} = 0$ and a two dimensional null-space, because the corresponding graph has two components.
- Let $N = |\mathcal{V}|$. The null space of $\mathbf{L}^{\text{disc}}$ is spanned by two vectors:
 - 1 one vector is $\mathbf{u}_N = \frac{1}{\sqrt{N}} \mathbf{1}$ (where $\frac{1}{\sqrt{N}}$ makes $\|\mathbf{u}_N\| = 1$),
 - 2 the other one which we can choose to be¹
 $\mathbf{u}_{N-1} \propto [\alpha \mathbf{1}^\top, \mathbf{1}^\top]^\top$ so that $\mathbf{u}_{N-1}^\top \mathbf{u}_N = 0$ and $\|\mathbf{u}_{N-1}\| = 1$.

¹ $\mathbf{a} \propto \mathbf{b}$ means $\mathbf{a}$ is proportional to $\mathbf{b}$, i.e. $\mathbf{a} = \gamma \mathbf{b}$



Null vectors of the Laplacian

- A similar partition can be done for the Laplacian:

$$\mathbf{L} = \mathbf{L}^{\text{disc}} + \delta \mathbf{L}$$

where $\delta \mathbf{L} = -\delta \mathbf{A}$, since $\mathbf{L} = \text{diag}(\mathbf{d}) - \mathbf{A}$, so the off diagonal elements of $\mathbf{L}$ have simply a $-$ sign compared to those of $\mathbf{A}$.

- We know that $\mathbf{L}^{\text{disc}}$ has Fiedler eigenvalue $\lambda_{N-1}^{\text{disc}} = 0$ and a two dimensional null-space, because the corresponding graph has two components.
- Let $N = |\mathcal{V}|$. The null space of $\mathbf{L}^{\text{disc}}$ is spanned by two vectors:
 - 1 one vector is $\mathbf{u}_N = \frac{1}{\sqrt{N}} \mathbf{1}$ (where $\frac{1}{\sqrt{N}}$ makes $\|\mathbf{u}_N\| = 1$),
 - 2 the other one which we can choose to be¹
 $\mathbf{u}_{N-1} \propto [\alpha \mathbf{1}^\top, \mathbf{1}^\top]^\top$ so that $\mathbf{u}_{N-1}^\top \mathbf{u}_N = 0$ and $\|\mathbf{u}_{N-1}\| = 1$.

¹ $\mathbf{a} \propto \mathbf{b}$ means $\mathbf{a}$ is proportional to $\mathbf{b}$, i.e. $\mathbf{a} = \gamma \mathbf{b}$



Null vectors of the Laplacian

- A similar partition can be done for the Laplacian:

$$\mathbf{L} = \mathbf{L}^{\text{disc}} + \delta \mathbf{L}$$

where $\delta \mathbf{L} = -\delta \mathbf{A}$, since $\mathbf{L} = \text{diag}(\mathbf{d}) - \mathbf{A}$, so the off diagonal elements of $\mathbf{L}$ have simply a $-$ sign compared to those of $\mathbf{A}$.

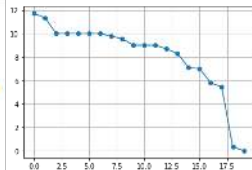
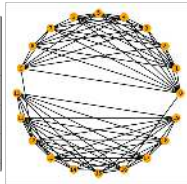
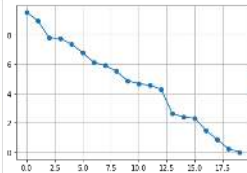
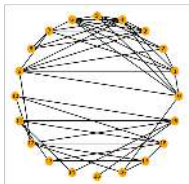
- We know that $\mathbf{L}^{\text{disc}}$ has Fiedler eigenvalue $\lambda_{N-1}^{\text{disc}} = 0$ and a two dimensional null-space, because the corresponding graph has two components.
- Let $N = |\mathcal{V}|$. The null space of $\mathbf{L}^{\text{disc}}$ is spanned by two vectors:
 - 1 one vector is $\mathbf{u}_N = \frac{1}{\sqrt{N}} \mathbf{1}$ (where $\frac{1}{\sqrt{N}}$ makes $\|\mathbf{u}_N\| = 1$),
 - 2 the other one which we can choose to be¹
 $\mathbf{u}_{N-1} \propto [\alpha \mathbf{1}^\top, \mathbf{1}^\top]^\top$ so that $\mathbf{u}_{N-1}^\top \mathbf{u}_N = 0$ and $\|\mathbf{u}_{N-1}\| = 1$.

¹ $\mathbf{a} \propto \mathbf{b}$ means $\mathbf{a}$ is proportional to $\mathbf{b}$, i.e. $\mathbf{a} = \gamma \mathbf{b}$



Graph with dense sections (clusters)

- We said an indicator of the presence of dense communities loosely connected is the **spectral gap** of the Laplacian $\gamma_k = |\lambda_{N-k-1} - \lambda_{N-k}|$ when there are k clusters in the graph
- Next we give a mathematical reason why.





Eigenvector of perturbed matrices

- Denote the eigenvectors of $\mathbf{L}$ by $\mathbf{u}_n$, those of $\mathbf{L}^{\text{disc}}$ by $\mathbf{u}_n^{\text{disc}}$ (we continue to focus on the case where there are two dense clusters loosely connected).

- Both $\mathbf{L}$ and $\mathbf{L}^{\text{disc}}$ are such that $\mathbf{L}\mathbf{1} = \mathbf{0}$ and $\mathbf{L}^{\text{disc}}\mathbf{1} = \mathbf{0}$, so $\mathbf{u}_N = \mathbf{u}_N^{\text{disc}} = \frac{1}{\sqrt{N}}\mathbf{1}$.

- The Fiedler eigenvalue $\lambda_{N-1} > 0$ but a theorem, by Weyl shows that $\lambda_{N-1} \leq \|\delta\mathbf{L}\|_2$

- For the eigenvectors, a result for perturbed matrices in general:

$$\mathbf{u}_{N-1}^\top \mathbf{u}_n^{\text{disc}} \leq \gamma^{-1} \|\delta\mathbf{L}\|_2, \quad \gamma = \min_k |\lambda_{N-k}(\mathbf{L}^{\text{disc}}) - \lambda_{N-k-1}(\mathbf{L}^{\text{disc}})|$$

is the minimum spectral gap.



Eigenvector of perturbed matrices

- Denote the eigenvectors of $\mathbf{L}$ by $\mathbf{u}_n$, those of $\mathbf{L}^{\text{disc}}$ by $\mathbf{u}_n^{\text{disc}}$ (we continue to focus on the case where there are two dense clusters loosely connected).
- Both $\mathbf{L}$ and $\mathbf{L}^{\text{disc}}$ are such that $\mathbf{L}\mathbf{1} = \mathbf{0}$ and $\mathbf{L}^{\text{disc}}\mathbf{1} = \mathbf{0}$, so $\mathbf{u}_N = \mathbf{u}_N^{\text{disc}} = \frac{1}{\sqrt{N}}\mathbf{1}$.
- The Fiedler eigenvalue $\lambda_{N-1} > 0$ but a theorem, by Weyl shows that $\lambda_{N-1} \leq \|\delta\mathbf{L}\|_2$
- For the eigenvectors, a result for perturbed matrices in general:

$$\mathbf{u}_{N-1}^\top \mathbf{u}_n^{\text{disc}} \leq \gamma^{-1} \|\delta\mathbf{L}\|_2, \quad \gamma = \min_k |\lambda_{N-k}(\mathbf{L}^{\text{disc}}) - \lambda_{N-k-1}(\mathbf{L}^{\text{disc}})|$$

is the minimum spectral gap.

Graph partitioning and communities



- A very common task in the analysis of graph data is to try to identify **clusters** that partition the network in communities.
- The most commonly used algorithm to do so is called **Spectral Clustering** and it works because of what we discussed.
- What we seek is a **partition** of a graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ that cuts the network in parts.
- We will see that the previous observations are helpful in providing intuition on why the algorithm works.

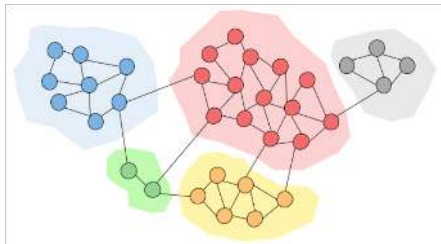


Graph partitioning and communities

Partition

- A partition of the nodes in the network, denoted by the subsets $\mathcal{V}_k \subseteq \mathcal{V}$ such that:

- 1) $\bigcup_{k=1}^K \mathcal{V}_k = \mathcal{V}$;
- 2) $\mathcal{V}_k \cap \mathcal{V}_{k'} = \emptyset$.





Graph partitioning

- We can consider the general case of weighted networks w_{uv}
- Let the complement of a subset $\mathcal{A}$ be denoted by $\overline{\mathcal{A}} = \mathcal{V}/\mathcal{A}$.
- All the measures of a *good partition* use the following definitions:

$$W(\mathcal{A}, \mathcal{B}) = \sum_{u \in \mathcal{A}} \sum_{v \in \mathcal{B}} w_{uv}$$

- When we partition the graph we want minimize the number (or weight) of the edges going outside $\mathcal{V}_k$, i.e. $W(\mathcal{V}_k, \overline{\mathcal{V}_k})$ as small as possible, but also have $\mathcal{V}_k$ with similar sizes.



Graph partitioning

- We can consider the general case of weighted networks w_{uv}
- Let the complement of a subset $\mathcal{A}$ be denoted by $\overline{\mathcal{A}} = \mathcal{V}/\mathcal{A}$.
- All the measures of a *good partition* use the following definitions:

$$W(\mathcal{A}, \mathcal{B}) = \sum_{u \in \mathcal{A}} \sum_{v \in \mathcal{B}} w_{uv}$$

- When we partition the graph we want minimize the number (or weight) of the edges going outside $\mathcal{V}_k$, i.e. $W(\mathcal{V}_k, \overline{\mathcal{V}_k})$ as small as possible, but also have $\mathcal{V}_k$ with similar sizes.



Graph partitioning

- We can consider the general case of weighted networks w_{uv}
- Let the complement of a subset $\mathcal{A}$ be denoted by $\overline{\mathcal{A}} = \mathcal{V}/\mathcal{A}$.
- All the measures of a *good partition* use the following definitions:

$$W(\mathcal{A}, \mathcal{B}) = \sum_{u \in \mathcal{A}} \sum_{v \in \mathcal{B}} w_{uv}$$

- When we partition the graph we want minimize the number (or weight) of the edges going outside $\mathcal{V}_k$, i.e. $W(\mathcal{V}_k, \overline{\mathcal{V}_k})$ as small as possible, but also have $\mathcal{V}_k$ with similar sizes.



Partitions and the Laplacian

Two facts allow us to write $W(\mathcal{V}_k, \overline{\mathcal{V}}_k)$ in terms of the Laplacian L :

- 1 The degree d_u of a node $u \in \mathcal{V}_k$ is the sum of the edges $(u, v), v \in \mathcal{V}_k$ and those with $v \in \overline{\mathcal{V}}_k$, i.e.

$$d_u = \sum_{v \in \mathcal{V}_k \setminus \{u\}} w_{u,v} + \sum_{v \in \overline{\mathcal{V}}_k} w_{u,v} ,$$

- 2 Recall that $L = \text{diag}(d) - A$. The diagonal elements of L are the degrees d_u and the off-diagonal are $-w_{uv}$.

$$\begin{aligned} W(\mathcal{V}_k, \overline{\mathcal{V}}_k) &= \sum_{u \in \mathcal{V}_k} \sum_{v \in \overline{\mathcal{V}}_k} w_{uv} \stackrel{(1)}{=} \sum_{u \in \mathcal{V}_k} \left(d_u - \sum_{v \in \mathcal{V}_k \setminus \{u\}} w_{u,v} \right) \\ &\stackrel{(2)}{=} \sum_{u \in \mathcal{V}_k} \sum_{v \in \mathcal{V}_k} [L]_{uv} \end{aligned}$$



Partitions and the Laplacian

Two facts allow us to write $W(\mathcal{V}_k, \overline{\mathcal{V}}_k)$ in terms of the Laplacian L :

- 1** The degree d_u of a node $u \in \mathcal{V}_k$ is the sum of the edges (u, v) , $v \in \mathcal{V}_k$ and those with $v \in \overline{\mathcal{V}}_k$, i.e.

$$d_u = \sum_{v \in \mathcal{V}_k \setminus \{u\}} w_{u,v} + \sum_{v \in \overline{\mathcal{V}}_k} w_{u,v} ,$$

- 2** Recall that $L = \text{diag}(d) - A$. The diagonal elements of L are the degrees d_u and the off-diagonal are $-w_{uv}$.

$$\begin{aligned} W(\mathcal{V}_k, \overline{\mathcal{V}}_k) &= \sum_{u \in \mathcal{V}_k} \sum_{v \in \overline{\mathcal{V}}_k} w_{uv} \stackrel{(1)}{=} \sum_{u \in \mathcal{V}_k} \left(d_u - \sum_{v \in \mathcal{V}_k \setminus \{u\}} w_{u,v} \right) \\ &\stackrel{(2)}{=} \sum_{u \in \mathcal{V}_k} \sum_{v \in \mathcal{V}_k} [L]_{uv} \end{aligned}$$



Partitions and the Laplacian

Two facts allow us to write $W(\mathcal{V}_k, \overline{\mathcal{V}}_k)$ in terms of the Laplacian $\mathbf{L}$:

- 1 The degree d_u of a node $u \in \mathcal{V}_k$ is the sum of the edges (u, v) , $v \in \mathcal{V}_k$ and those with $v \in \overline{\mathcal{V}}_k$, i.e.

$$d_u = \sum_{v \in \mathcal{V}_k \setminus \{u\}} w_{u,v} + \sum_{v \in \overline{\mathcal{V}}_k} w_{u,v} ,$$

- 2 Recall that $\mathbf{L} = \text{diag}(\mathbf{d}) - \mathbf{A}$. The diagonal elements of $\mathbf{L}$ are the degrees d_u and the off-diagonal are $-w_{uv}$.

$$\begin{aligned} W(\mathcal{V}_k, \overline{\mathcal{V}}_k) &= \sum_{u \in \mathcal{V}_k} \sum_{v \in \overline{\mathcal{V}}_k} w_{uv} \stackrel{(1)}{=} \sum_{u \in \mathcal{V}_k} \left(d_u - \sum_{v \in \mathcal{V}_k \setminus \{u\}} w_{u,v} \right) \\ &\stackrel{(2)}{=} \sum_{u \in \mathcal{V}_k} \sum_{v \in \mathcal{V}_k} [\mathbf{L}]_{uv} \end{aligned}$$



Partitions and the Laplacian

Two facts allow us to write $W(\mathcal{V}_k, \overline{\mathcal{V}}_k)$ in terms of the Laplacian L :

- 1** The degree d_u of a node $u \in \mathcal{V}_k$ is the sum of the edges (u, v) , $v \in \mathcal{V}_k$ and those with $v \in \overline{\mathcal{V}}_k$, i.e.

$$d_u = \sum_{v \in \mathcal{V}_k \setminus \{u\}} w_{u,v} + \sum_{v \in \overline{\mathcal{V}}_k} w_{u,v} ,$$

- 2** Recall that $L = \text{diag}(\mathbf{d}) - \mathbf{A}$. The diagonal elements of L are the degrees d_u and the off-diagonal are $-w_{uv}$.

$$\begin{aligned} W(\mathcal{V}_k, \overline{\mathcal{V}}_k) &= \sum_{u \in \mathcal{V}_k} \sum_{v \in \overline{\mathcal{V}}_k} w_{uv} \stackrel{(1)}{=} \sum_{u \in \mathcal{V}_k} \left(d_u - \sum_{v \in \mathcal{V}_k \setminus \{u\}} w_{u,v} \right) \\ &\stackrel{(2)}{=} \sum_{u \in \mathcal{V}_k} \sum_{v \in \mathcal{V}_k} [L]_{uv} \end{aligned}$$



Balanced graph clustering

Balanced graph clustering times at identifying **communities** by *minimizing* the following metric:

$$\text{RatioCut}(\mathcal{V}_1, \dots, \mathcal{V}_K) := \sum_{k=1}^K \frac{1}{|\mathcal{V}_k|} W(\mathcal{V}_k, \bar{\mathcal{V}}_k)$$

- We can define the following vectors that are associated with a partition $\mathcal{V}_1, \dots, \mathcal{V}_K$:

$$\mathbf{p}_k \text{ such that } [\mathbf{p}_k]_u = \begin{cases} \frac{1}{\sqrt{|\mathcal{V}_k|}} & u \in \mathcal{V}_k \\ 0 & \text{else} \end{cases} \quad k = 1, \dots, K$$

- Then, we can write

$$W(\mathcal{V}_k, \bar{\mathcal{V}}_k) = \sum_{u \in \mathcal{V}_k} \sum_{v \in \mathcal{V}_k} [L]_{uv} = |\mathcal{V}_k| \mathbf{p}_k^\top L \mathbf{p}_k$$



Balanced graph clustering

Balanced graph clustering times at identifying **communities** by *minimizing* the following metric:

$$\text{RatioCut}(\mathcal{V}_1, \dots, \mathcal{V}_K) := \sum_{k=1}^K \frac{1}{|\mathcal{V}_k|} W(\mathcal{V}_k, \bar{\mathcal{V}}_k)$$

- We can define the following vectors that are associated with a partition $\mathcal{V}_1, \dots, \mathcal{V}_K$:

$$\mathbf{p}_k \text{ such that } [\mathbf{p}_k]_u = \begin{cases} \frac{1}{\sqrt{|\mathcal{V}_k|}} & u \in \mathcal{V}_k \\ 0 & \text{else} \end{cases} \quad k = 1, \dots, K$$

- Then, we can write

$$W(\mathcal{V}_k, \bar{\mathcal{V}}_k) = \sum_{u \in \mathcal{V}_k} \sum_{v \in \mathcal{V}_k} [\mathbf{L}]_{uv} = |\mathcal{V}_k| \mathbf{p}_k^\top \mathbf{L} \mathbf{p}_k$$



Balanced graph clustering

- Note that $\|\mathbf{p}_k\|^2 = \mathbf{p}_k^\top \mathbf{p}_k = 1$, and for $m \neq k$, $\mathbf{p}_k^\top \mathbf{p}_m = 0$.
- Introducing the matrix $\mathbf{P} := (\mathbf{p}_1, \dots, \mathbf{p}_K)$ we note that

$$\mathbf{p}_k^\top \mathbf{p}_m = \begin{cases} 1 & k = m \\ 0 & \text{else} \end{cases} \Rightarrow \mathbf{P}^\top \mathbf{P} = \mathbf{I}$$

- In addition we can write the ratio cut as follows:

$$\begin{aligned} \text{RatioCut}(\mathcal{V}_1, \dots, \mathcal{V}_K) &= \sum_{k=1}^K \frac{1}{|\mathcal{V}_k|} \overbrace{\sum_{u \in \mathcal{V}_k} \sum_{v \in \mathcal{V}_k} [\mathbf{L}]_{uv}}^{|\mathcal{V}_k| \mathbf{p}_k^\top \mathbf{L} \mathbf{p}_k} \\ &= \sum_{k=1}^K \mathbf{p}_k^\top \mathbf{L} \mathbf{p}_k = \text{Tr}(\mathbf{P}^\top \mathbf{L} \mathbf{P}) \end{aligned}$$

where $\text{Tr}(\mathbf{A}) := \sum_{v=1} \mathbf{A}_{vv}$ is the **trace of matrix $\mathbf{A}$**

- Note that $\text{Tr}(\mathbf{A}) = \text{Tr}(\mathbf{U}_A \mathbf{\Lambda}_A \mathbf{U}_A^\top) = \text{Tr}(\mathbf{\Lambda}_A) \equiv \sum_{v=1} \lambda_v(\mathbf{A})$, where $\lambda_v(\mathbf{A})$ are the eigenvalues of $\mathbf{A}$.



Minimum Ratio Cut criterion

Minimum Ratio Cut clustering

The optimum partition using the RatioCut metric

$$\min_{\mathcal{V}_1, \dots, \mathcal{V}_K} \text{RatioCut}(\mathcal{V}_1, \dots, \mathcal{V}_K)$$

can be formulated in terms of the Laplacian as follows:

$$\min_{\mathbf{P}} \text{Tr}(\mathbf{P}^\top \mathbf{L} \mathbf{P}), \quad \text{subject to } \mathbf{P}^\top \mathbf{P} = \mathbf{I}, \quad \mathbf{P} \in \mathcal{P}$$

where $\mathcal{P}$ is the set of $N \times K$ matrices $\mathbf{P}$ whose entries are $[\mathbf{P}]_{uk} = p_{uk} \in \{0, \alpha_k\}$, $\alpha_k^2 \|\mathbf{p}_k\|_0 = 1$



Unnormalized spectral clustering: min RatioCut

- The min RatioCut problem is a **combinatorial problem** because of the this constraint expressed by $P \in \mathcal{P}$ that requires searching for the right support of each vector p_k , $k = 1, \dots, K$
- Consider a constrained optimization:

$$\min_x C(x) \quad \text{subject to } x \in \mathcal{X}$$

- **Relaxation:** Relaxing the optimization problem means loosening either completely or in part the constraint $x \in \mathcal{X}$ by finding an easier constraint to deal with.



Unnormalized spectral clustering: min RatioCut

- The min RatioCut problem is a **combinatorial problem** because of the this constraint expressed by $P \in \mathcal{P}$ that requires searching for the right support of each vector p_k , $k = 1, \dots, K$
- Consider a constrained optimization:

$$\min_x C(x) \quad \text{subject to } x \in \mathcal{X}$$

- **Relaxation:** Relaxing the optimization problem means loosening either completely or in part the constraint $x \in \mathcal{X}$ by finding an easier constraint to deal with.

Unnormalized spectral clustering: min RatioCut



- The min RatioCut problem is a **combinatorial problem** because of the this constraint expressed by $P \in \mathcal{P}$ that requires searching for the right support of each vector p_k , $k = 1, \dots, K$
- Consider a constrained optimization:

$$\min_x C(x) \quad \text{subject to } x \in \mathcal{X}$$

- **Relaxation:** Relaxing the optimization problem means loosening either completely or in part the constraint $x \in \mathcal{X}$ by finding an easier constraint to deal with.



Relaxation

- More formally, the relaxed set $\mathcal{X}'$ is such that $\mathcal{X} \subset \mathcal{X}'$.

$$\min_x C(x) \quad \text{subject to } x \in \mathcal{X}'$$

- If $\mathcal{X}' \equiv \mathbb{R}^N$ the constraint is completely relaxed, i.e. becomes unconstrained.
- It holds that the minimum objective of the relaxed problem has always a minimum cost that is no larger than the minimum cost of the original problem.
- So if x_r is the solution of the relaxed problem and x_o that of the original problem can achieve a lower cost (it includes all previous possible values and more)

$$\min_{x \in \mathcal{X}'} C(x) \leq \min_{x \in \mathcal{X}} C(x).$$



Relaxation

- More formally, the relaxed set $\mathcal{X}'$ is such that $\mathcal{X} \subset \mathcal{X}'$.

$$\min_x C(\mathbf{x}) \quad \text{subject to} \quad \mathbf{x} \in \mathcal{X}'$$

- If $\mathcal{X}' \equiv \mathbb{R}^N$ the constraint is completely relaxed, i.e. becomes unconstrained.
- It holds that the minimum objective of the relaxed problem has always a minimum cost that is no larger than the minimum cost of the original problem.
- So if $\mathbf{x}_r$ is the solution of the relaxed problem and $\mathbf{x}_o$ that of the original problem can achieve a lower cost (it includes all previous possible values and more)

$$\min_{\mathbf{x} \in \mathcal{X}'} C(\mathbf{x}) \leq \min_{\mathbf{x} \in \mathcal{X}} C(\mathbf{x}).$$



Min RatioCut relaxation

- The spectral clustering algorithm comes from relaxing the constraint $\mathbf{P} \in \mathcal{P}$. In other words, by solving instead

$$\min_{\mathbf{P}} \text{Tr}(\mathbf{P}^{\top} \mathbf{L} \mathbf{P}), \quad \text{subject to } \mathbf{P}^{\top} \mathbf{P} = \mathbf{I}.$$

where $\mathbf{P}^{\top} \mathbf{P} = \mathbf{I}$ means that $\mathbf{P}$ is simply orthonormal, does not have to be necessarily a diagonally scaled permutation matrix.

- This problem is also non convex, because it of the orthogonality constraint, but has a solution:

$$\mathbf{P}_{\text{opt}} = \mathbf{U}_{\text{min}}$$

where the columns of $\mathbf{U}_{\text{min}}$ are the K least significant eigenvectors of $\mathbf{L}$.



Min RatioCut relaxation

- The spectral clustering algorithm comes from relaxing the constraint $\mathbf{P} \in \mathcal{P}$. In other words, by solving instead

$$\min_{\mathbf{P}} \text{Tr}(\mathbf{P}^\top \mathbf{L} \mathbf{P}), \quad \text{subject to } \mathbf{P}^\top \mathbf{P} = \mathbf{I}.$$

where $\mathbf{P}^\top \mathbf{P} = \mathbf{I}$ means that $\mathbf{P}$ is simply orthonormal, does not have to be necessarily a diagonally scaled permutation matrix.

- This problem is also non convex, because it of the orthogonality constraint, but has a solution:

$$\mathbf{P}_{\text{opt}} = \mathbf{U}_{\text{min}}$$

where the columns of $\mathbf{U}_{\text{min}}$ are the K least significant eigenvectors of $\mathbf{L}$.



Minimum RatioCut bound

- With $P_{\text{opt}} = U_{\text{min}}$ the minimum objective becomes

$$\min_P \text{Tr}(U_{\text{min}}^T L U_{\text{min}}) = \text{Tr}(\Lambda_{\text{min}})$$

where Λ_{min} is the diagonal matrix of

$\lambda_{N-K} \geq \dots \geq \lambda_{N-1} \geq \lambda_N = 0$ are the smallest K eigenvalues of the Laplacian matrix L .

- Since the relaxation results in a minimum that is no larger than original problem (i.e. $\min_{x \in \mathcal{X}'} C(x) \leq \min_{x \in \mathcal{X}} C(x)$):

$$\min_{\mathcal{V}_1, \dots, \mathcal{V}_K} \text{RatioCut}(\mathcal{V}_1, \dots, \mathcal{V}_K) \geq \sum_{k=1}^{K-1} \lambda_{N-k}$$

- Now, how do we go from U_{min} to the partition $(\mathcal{V}_1, \dots, \mathcal{V}_K)$?



Minimum RatioCut bound

- With $P_{\text{opt}} = U_{\min}$ the minimum objective becomes

$$\min_P \text{Tr}(U_{\min}^{\top} L U_{\min}) = \text{Tr}(\Lambda_{\min})$$

where $\Lambda_{\min}$ is the diagonal matrix of

$\lambda_{N-K} \geq \dots \geq \lambda_{N-1} \geq \lambda_N = 0$ are the smallest K eigenvalues of the Laplacian matrix L .

- Since the relaxation results in a minimum that is no larger than original problem (i.e. $\min_{x \in \mathcal{X}'} C(x) \leq \min_{x \in \mathcal{X}} C(x)$):

$$\min_{\mathcal{V}_1, \dots, \mathcal{V}_K} \text{RatioCut}(\mathcal{V}_1, \dots, \mathcal{V}_K) \geq \sum_{k=1}^{K-1} \lambda_{N-k}$$

- Now, how do we go from $U_{\min}$ to the partition $(\mathcal{V}_1, \dots, \mathcal{V}_K)$?



Minimum RatioCut bound

- With $P_{\text{opt}} = U_{\min}$ the minimum objective becomes

$$\min_P \text{Tr}(U_{\min}^T L U_{\min}) = \text{Tr}(\Lambda_{\min})$$

where $\Lambda_{\min}$ is the diagonal matrix of

$\lambda_{N-K} \geq \dots \geq \lambda_{N-1} \geq \lambda_N = 0$ are the smallest K eigenvalues of the Laplacian matrix L .

- Since the relaxation results in a minimum that is no larger than original problem (i.e. $\min_{x \in \mathcal{X}'} C(x) \leq \min_{x \in \mathcal{X}} C(x)$):

$$\min_{\mathcal{V}_1, \dots, \mathcal{V}_K} \text{RatioCut}(\mathcal{V}_1, \dots, \mathcal{V}_K) \geq \sum_{k=1}^{K-1} \lambda_{N-k}$$

- Now, how do we go from $U_{\min}$ to the partition $(\mathcal{V}_1, \dots, \mathcal{V}_K)$?



Unnormalized spectral clustering

- The most common approach is to use the K – **means** algorithm to the **rows** of $U_{\min}$, to now cluster the entries of the vectors.
- Why on earth would that work?
- First we need to understand what K –means does

K-means

K-means partition N observations into K clusters in which each observation belongs to the cluster with the nearest mean.



Unnormalized spectral clustering

- The most common approach is to use the K – **means** algorithm to the **rows** of $U_{\min}$, to now cluster the entries of the vectors.
- Why on earth would that work?
 - First we need to understand what K –means does

K-means

K-means partition N observations into K clusters in which each observation belongs to the cluster with the nearest mean.



Unnormalized spectral clustering

- The most common approach is to use the K – **means** algorithm to the **rows** of $U_{\min}$, to now cluster the entries of the vectors.
- Why on earth would that work?
- First we need to understand what K –means does

K-means

K-means partition N observations into K clusters in which each observation belongs to the cluster with the nearest mean.



K-means algorithm

- Given a set of points $\mathcal{X} = \mathbf{x}_1, \dots$ K -means seeks a partition of $\mathcal{X}$, in K disjoint subsets of points $\{\mathcal{X}_1, \dots, \mathcal{X}_K\}$ that targets the following objective:

$$\min_{\{\mathcal{X}_1, \dots, \mathcal{X}_K\}} \sum_{k=1}^K \sum_{i: \mathbf{x}_i \in \mathcal{X}_k} \|\mathbf{x}_i - \boldsymbol{\mu}_k\|^2, \quad \text{where } \boldsymbol{\mu}_k = \frac{1}{|\mathcal{X}_k|} \sum_{i: \mathbf{x}_i \in \mathcal{X}_k} \mathbf{x}_i$$

where $\boldsymbol{\mu}_k$ are called the centroids.



K-means algorithm on $U_{\min}$

- Let us denote the rows of $U_{\min}$ as follows

$$U_{\min} = \begin{bmatrix} \mathbf{u}_1^{\text{row}} \\ \vdots \\ \mathbf{u}_N^{\text{row}} \end{bmatrix}$$

we just apply K -means to $\mathbf{u}_n^{\text{row}}, n = 1, \dots, N$ as the set of points to cluster.

- The if $\mathbf{u}_n^{\text{row}}, n = 1, \dots, N$ is in cluster k then node n is in the graph cluster $\mathcal{V}_k$



K-means algorithm on $U_{\min}$

- Let us denote the rows of $U_{\min}$ as follows

$$U_{\min} = \begin{bmatrix} \mathbf{u}_1^{\text{row}} \\ \vdots \\ \mathbf{u}_N^{\text{row}} \end{bmatrix}$$

we just apply K -means to $\mathbf{u}_n^{\text{row}}, n = 1, \dots, N$ as the set of points to cluster.

- The if $\mathbf{u}_n^{\text{row}}, n = 1, \dots, N$ is in cluster k then node n is in the graph cluster $\mathcal{V}_k$

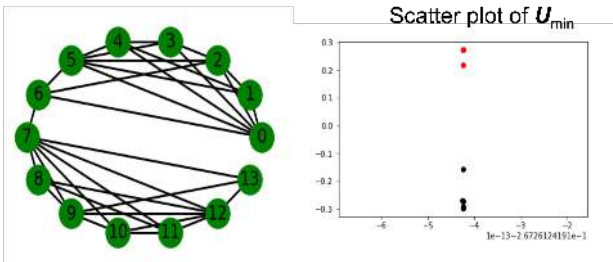


The intuition

- We already discussed this, we said that impact L can be represented as follows

$$L = \begin{bmatrix} L_1 & \mathbf{0} \\ \mathbf{0} & L_2 \end{bmatrix} + \delta L, \quad \|\delta L\| \ll \|L\|$$

- We already had marked in red and black the points that correspond to the different clusters and said that the same can be generalized for a graph that has $K > 2$ clusters.

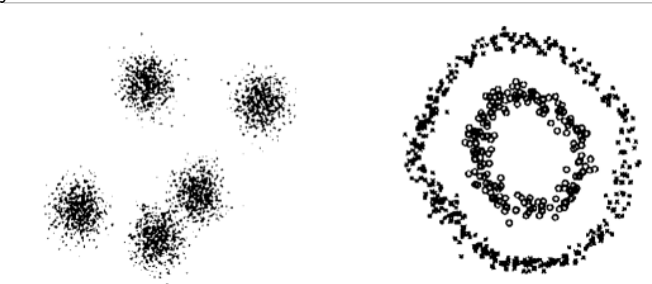


Normalized laplacian matrix of weighted graphs



Enschede University

- Normalized Laplacians are also used to understand the topological structure of data and identify if they can be partitioned.
- In both figures the points scattered are two dimensional real vectors x_i and clustering algorithms are *unsupervised* method to identify the indexes of these sets

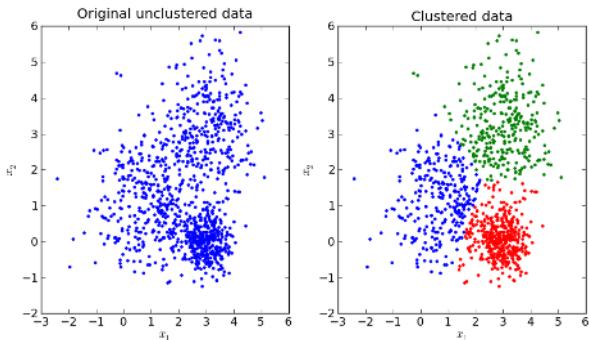


- Spectral clustering is much more effective for data scattered in groups around non convex shapes, like the ones on the right figure.

K-Means vs. Spectral Clustering



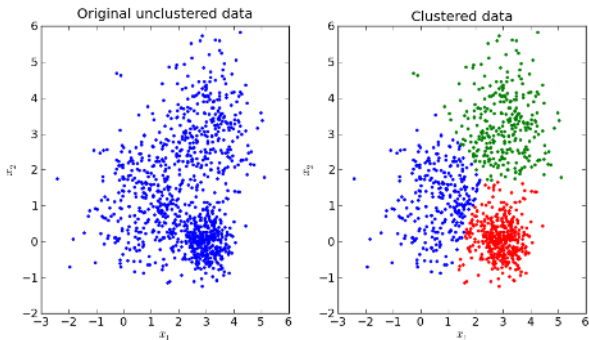
- The left case happens to be a **mixture of Gaussian**: You have a categorical distribution $P(k) : \sum_{k=1}^K P(k) = 1$ and with probability $P(k)$ you draw a sample out of $x \sim \mathcal{N}(\mu_k, R_k)$.
- For the points obtained, you can apply K-means directly to the points x_i



K-Means vs. Spectral Clustering



- The left case happens to be a **mixture of Gaussian**: You have a categorical distribution $P(k) : \sum_{k=1}^K P(k) = 1$ and with probability $P(k)$ you draw a sample out of $\mathbf{x} \sim \mathcal{N}(\boldsymbol{\mu}_k, \mathbf{R}_k)$.
- For the points obtained, you can apply K-means directly to the points $\mathbf{x}_i$

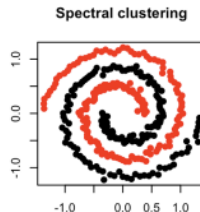
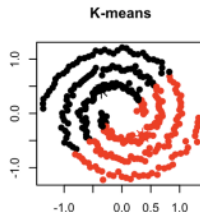
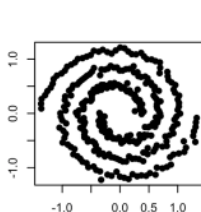


K-Means vs. Spectral Clustering



Cornell University

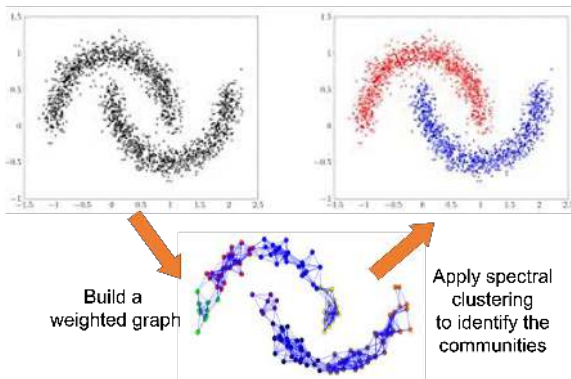
- If you apply K-means directly to the points x_i for the other case you do not separate correctly the data



Normalized laplacian matrix of weighted graphs



- What is done is to first define weights that rapidly monotonically *decrease* with the data points distance, denoting weak affinity. You threshold those that are too small.
- Non convex clusters of points will form tight nit communities





Normalized laplacian matrix of weighted graphs

- Let the data points be $\mathbf{x}_i, i = 1, \dots, n$, and $d(\mathbf{x}_i, \mathbf{x}_j)$ a measure of their distance, affinity matrices (e.g. for points $\mathbf{x}_i \in \mathbb{R}^D$ we can use the Frobenious norm $\|\mathbf{x}_i - \mathbf{x}_j\|$). One option often used is the so called Gaussian kernel:

$$[\mathbf{A}]_{ij} = w_{ij} = e^{-\frac{\|\mathbf{x}_i - \mathbf{x}_j\|^2}{2}} \Rightarrow \mathbf{L} = \text{diag}(\mathbf{d}) - \mathbf{A}$$

- The weights w_{ij} measure of *affinity* and can be used to define a weighted adjacency matrix.
- Instead of applying K-means to the points $\mathbf{x}_i$ spectral clustering can be applied to separate these points, by applying K-means to the K least significant eigenvectors $\mathbf{U}_{\min}$ of this Laplacian.

Normalized laplacian matrix of weighted graphs



- Let the data points be $\mathbf{x}_i, i = 1, \dots, n$, and $d(\mathbf{x}_i, \mathbf{x}_j)$ a measure of their distance, affinity matrices (e.g. for points $\mathbf{x}_i \in \mathbb{R}^D$ we can use the Frobenious norm $\|\mathbf{x}_i - \mathbf{x}_j\|$). One option often used is the so called Gaussian kernel:

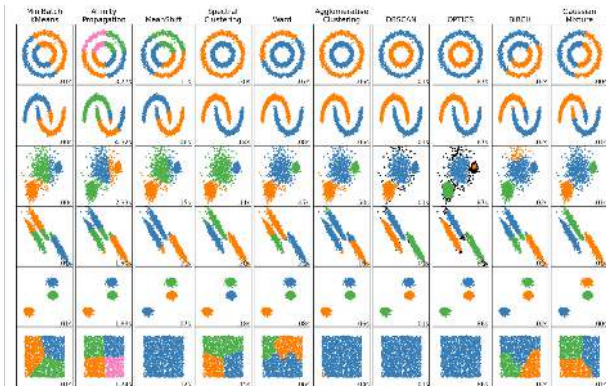
$$[\mathbf{A}]_{ij} = w_{ij} = e^{-\frac{\|\mathbf{x}_i - \mathbf{x}_j\|^2}{2}} \Rightarrow \mathbf{L} = \text{diag}(\mathbf{d}) - \mathbf{A}$$

- The weights w_{ij} measure of *affinity* and can be used to define a weighted adjacency matrix.
- Instead of applying K-means to the points $\mathbf{x}_i$ spectral clustering can be applied to separate these points, by applying K-means to the K least significant eigenvectors $\mathbf{U}_{\min}$ of this Laplacian.



Clustering methods

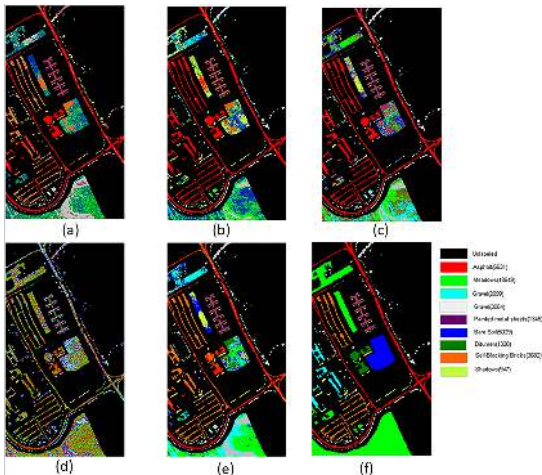
- Of course K -means is also a clustering method and there are many variations, but spectral clustering is very popular and effective.





Applications

- There are many interesting applications

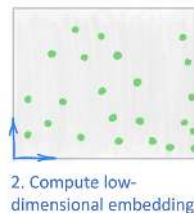
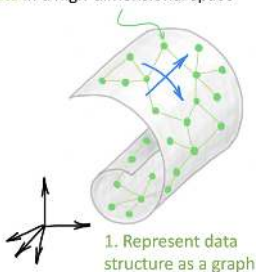




Applications of spectral clustering

- For many data, not necessarily 2-D, that are expected to have a low dimensional structure that is non convex, spectral clustering is the go-to method for - unsupervised classification

intrinsically low-dimensional
data in a high-dimensional space



- In some cases, one seeks to map them in a lower dimensions first through what is called an embedding



Outline

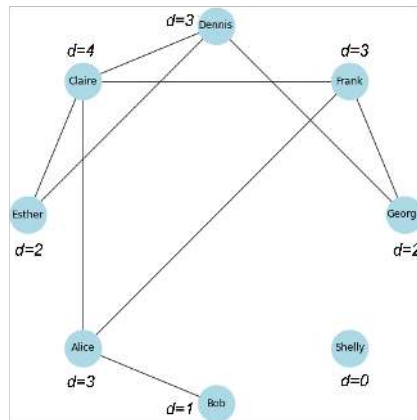
- 1 Introduction
- 2 Betti number and Hodge Laplacian
- 3 Graph partitioning
- 4 The ranking of nodes: centrality measures**
- 5 Conclusions



Nodal degree

Feature extraction

- We already defined the degree $d_v = |\mathcal{N}_v|$ is the number of neighbors of a node. The degree sequence is a feature of the graph.
- For a directed graph we have in degree and out degrees d_v^+ and d_v^- , which provide for each node a 2D feature component.
- For simplicity we will focus on undirected graphs.





Degree distribution

Feature extraction

- For the whole graph the empirical distribution of the nodal degree is a feature vector: it is the fraction of nodes that has a given nodal degree. Mathematically:

$$p_k = \hat{P}(d_v = k) = \frac{1}{|\mathcal{V}|} \sum_{v \in \mathcal{V}} \delta(d_v - k), \quad k = 1, \dots, \bar{d}$$

where $\bar{d} = \max_{u \in \mathcal{V}}(d_u)$ and $\delta(k)$ is the Kronecker delta, i.e the indicator function defined as:

$$\delta(k) := \begin{cases} 1 & \text{for } k = 0 \\ 0 & \text{else} \end{cases}$$

- Note that the average degree:

$$d_{av} = \frac{1}{|\mathcal{V}|} \sum_{v \in \mathcal{V}} d_v = \sum_{k=1}^{\bar{d}} k p_k$$



Degree distribution

Feature extraction

- For the whole graph the empirical distribution of the nodal degree is a feature vector: it is the fraction of nodes that has a given nodal degree. Mathematically:

$$p_k = \hat{P}(d_v = k) = \frac{1}{|\mathcal{V}|} \sum_{v \in \mathcal{V}} \delta(d_v - k), \quad k = 1, \dots, \bar{d}$$

where $\bar{d} = \max_{u \in \mathcal{V}}(d_u)$ and $\delta(k)$ is the Kronecker delta, i.e the indicator function defined as:

$$\delta(k) := \begin{cases} 1 & \text{for } k = 0 \\ 0 & \text{else} \end{cases}$$

- Note that the average degree:

$$d_{\text{av}} = \frac{1}{|\mathcal{V}|} \sum_{v \in \mathcal{V}} d_v = \sum_{k=1}^{\bar{d}} k p_k$$



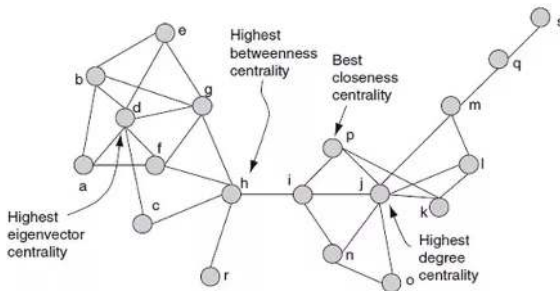
Node centrality measures

graph analysis

Centrality

This measure is used to establish the **most important node in the network**

- They find different nodes, since the form of *importance* is different and has a functional meaning





Degree centrality

graph analysis

- Degree centrality is one measure. The central node is the one that has maximum degree, i.e.

$$v^{\text{deg}} = \operatorname{argmax}_v d_v.$$

- Is very simple but often very indicative of importance in social networks and of congestion in infrastructures

Degree centrality

graph analysis



- Degree centrality is one measure. The central node is the one that has maximum degree, i.e.

$$v^{\text{deg}} = \operatorname{argmax}_v d_v.$$

- Is very simple but often very indicative of importance in social networks and of congestion in infrastructures



Eigenvector centrality

graph analysis

- Another popular measure is called **Eigenvector Centrality** that is defined with a recurrence equation, which captures the fact that an central node must have **central neighbors**:

$$z_u \propto \sum_{v \in \mathcal{N}_u} z_v \quad u \in \mathcal{V} \quad \Rightarrow \quad \lambda \mathbf{z} = \mathbf{A} \mathbf{z}$$

- Which eigenvector would we use, given that the score is proportional to λ and there are N of them?



Eigenvector centrality

graph analysis

- Another popular measure is called **Eigenvector Centrality** that is defined with a recurrence equation, which captures the fact that a central node must have **central neighbors**:

$$z_u \propto \sum_{v \in \mathcal{N}_u} z_v \quad u \in \mathcal{V} \quad \Rightarrow \quad \lambda \mathbf{z} = \mathbf{A} \mathbf{z}$$

- Which eigenvector would we use, given that the score is proportional to λ and there are N of them?



Eigenvector centrality

graph analysis

- Note that A has non negative entries: **Perron Frobenious Theorem** guarantees that the **maximum eigenvector has positive entries**.
- Hence the eigenvector centrality of a node a node is the corresponding entry of the **leading eigenvector**, which will assign a **positive score to every node**
- Let z_1 be the leading eigenvector, i.e. $\lambda_1(A)z_1 = Az_1$. The most central node is²:

$$v^{\text{eig}} = \operatorname{argmax}_v [z_1]_u$$

- Again, for disconnected networks is more meaningful to have a central node per component.

²Any scaling of the leading eigenvector will give the same result.



Eigenvector centrality

graph analysis

- Note that $\mathbf{A}$ has non negative entries: **Perron Frobenious Theorem** guarantees that the **maximum eigenvector has positive entries**.
- Hence the eigenvector centrality of a node a node is the corresponding entry of the **leading eigenvector**, which will assign a **positive score to every node**
- Let $\mathbf{z}_1$ be the leading eigenvector, i.e. $\lambda_1(\mathbf{A})\mathbf{z}_1 = \mathbf{A}\mathbf{z}_1$. The most central node is²:

$$v^{\text{eig}} = \operatorname{argmax}_v [\mathbf{z}_1]_u$$

- Again, for disconnected networks is more meaningful to have a central node per component.

²Any scaling of the leading eigenvector will give the same result.



Eigenvector centrality

graph analysis

- Note that $\mathbf{A}$ has non negative entries: **Perron Frobenious Theorem** guarantees that the **maximum eigenvector has positive entries**.
- Hence the eigenvector centrality of a node a node is the corresponding entry of the **leading eigenvector**, which will assign a **positive score to every node**
- Let $\mathbf{z}_1$ be the leading eigenvector, i.e. $\lambda_1(\mathbf{A})\mathbf{z}_1 = \mathbf{A}\mathbf{z}_1$. The most central node is²:

$$v^{\text{eig}} = \operatorname{argmax}_v [\mathbf{z}_1]_u$$

- Again, for disconnected networks is more meaningful to have a central node per component.

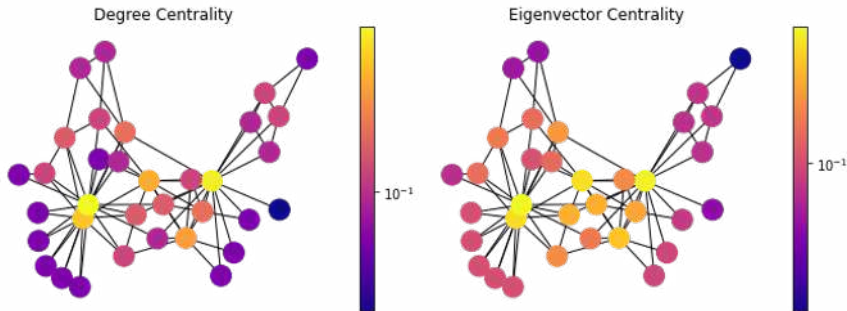
²Any scaling of the leading eigenvector will give the same result.



Eigenvector centrality example

graph analysis

- The example of the Zachary's Karate Club for both measures illustrates similarities and differences between the two metrics





Outline

- 1 Introduction
- 2 Betti number and Hodge Laplacian
- 3 Graph partitioning
- 4 The ranking of nodes: centrality measures
- 5 Conclusions**



Conclusions

- We studied in detail the problem of graph partitioning and saw how the Laplacian plays a central role in it
- We introduced some metrics to characterize them and compare nodes.
- In the next lecture we will continue to refine the notion of centrality, and introduce the general notion of graph embedding and introduce other graph features that are used to identify structural characteristics.