



Lecture 20 - Unsupervised GSP Algorithms

ECE 5260 – ORIE 5735

Anna Scaglione

Cornell Tech

April 20, 2026

Content



- 1 Review GSP
- 2 Unsupervised learning tasks
- 3 Low Rank Property of graph signals
 - Sampling and interpolation theorem
 - From graph data to graph topology

Outline



- 1 Review GSP
- 2 Unsupervised learning tasks
- 3 Low Rank Property of graph signals

Graph Filters



- Graph Filters are **shift invariant** linear operators, i.e.
 $\mathbf{S}\mathcal{H}(\mathbf{S}) \equiv \mathcal{H}(\mathbf{S})\mathbf{S}$. To be invariant w.r.t. the GSO a **graph filter** must be a matrix polynomial:

$$\mathcal{H}(\mathbf{S}) \triangleq \sum_{\ell=-\infty}^{+\infty} h_{\ell} \mathbf{S}^{\ell}. \quad (1)$$

- Examples: finite order $h_0 + h_1 \mathbf{S} + h_2 \mathbf{S}^2$, of infinite order filters:
 $\mathcal{H}(\mathbf{S}) = (\mathbb{I} - \alpha \mathbf{S})^{-1}$, $\mathcal{H}(\mathbf{S}) = e^{\alpha \mathbf{S}}$.
- Graph/network data are **filtered** graph signals, i.e.,

$$\mathbf{x} = \mathcal{H}(\mathbf{S})\mathbf{s}. \quad (2)$$

- The signal/observation is $\mathbf{x}$ while $\mathbf{s}$ is viewed as the **excitation**.

Graph Filters



- Graph Filters are **shift invariant** linear operators, i.e.
 $\mathbf{S}\mathcal{H}(\mathbf{S}) \equiv \mathcal{H}(\mathbf{S})\mathbf{S}$. To be invariant w.r.t. the GSO a **graph filter** must be a matrix polynomial:

$$\mathcal{H}(\mathbf{S}) \triangleq \sum_{\ell=-\infty}^{+\infty} h_{\ell} \mathbf{S}^{\ell}. \quad (1)$$

- Examples: finite order $h_0 + h_1 \mathbf{S} + h_2 \mathbf{S}^2$, of infinite order filters:
 $\mathcal{H}(\mathbf{S}) = (\mathbb{I} - \alpha \mathbf{S})^{-1}$, $\mathcal{H}(\mathbf{S}) = e^{\alpha \mathbf{S}}$.
- Graph/network data are **filtered** graph signals, i.e.,

$$\mathbf{x} = \mathcal{H}(\mathbf{S})\mathbf{s}. \quad (2)$$

- The signal/observation is $\mathbf{x}$ while $\mathbf{s}$ is viewed as the **excitation**.

Graph Filters



- Graph Filters are **shift invariant** linear operators, i.e.
 $\mathbf{S}\mathcal{H}(\mathbf{S}) \equiv \mathcal{H}(\mathbf{S})\mathbf{S}$. To be invariant w.r.t. the GSO a **graph filter** must be a matrix polynomial:

$$\mathcal{H}(\mathbf{S}) \triangleq \sum_{\ell=-\infty}^{+\infty} h_{\ell} \mathbf{S}^{\ell}. \quad (1)$$

- Examples: finite order $h_0 + h_1 \mathbf{S} + h_2 \mathbf{S}^2$, of infinite order filters:
 $\mathcal{H}(\mathbf{S}) = (\mathbb{I} - \alpha \mathbf{S})^{-1}$, $\mathcal{H}(\mathbf{S}) = e^{\alpha \mathbf{S}}$.
- Graph/network data are **filtered** graph signals, i.e.,

$$\mathbf{x} = \mathcal{H}(\mathbf{S})\mathbf{s}. \quad (2)$$

- The signal/observation is $\mathbf{x}$ while $\mathbf{s}$ is viewed as the **excitation**.

Graph Fourier Transform (GFT)



- The Graph Fourier Transform (GFT) basis are the eigen-vectors of the GSO (generalization of the DFT case for the cycle graph).
- $\mathbf{S} = \mathbf{U}\mathbf{\Lambda}\mathbf{U}^{-1}$. The columns of $\mathbf{U}$ are the signals that are **invariant** w.r.t. the GSO $\mathbf{S}$

$$\tilde{\mathbf{x}} = \overbrace{\mathbf{U}^{-1}}^{\text{GFT}} \mathbf{x} \quad (3)$$

and the eigenvalues are the graph *frequencies* from low to high:

$$0 = \lambda_1 \leq \dots \leq \lambda_N.$$

- The **transfer function** of the GF is:

$$\tilde{\mathcal{H}}(\mathbf{\Lambda}) = \mathbf{diag}(\tilde{\mathbf{h}}) \quad \text{where} \quad [\tilde{\mathbf{h}}]_i = \tilde{h}_i = \sum_{\ell} h_{\ell} \lambda_i^{\ell}$$

node coordinates: $\mathbf{x} = \mathcal{H}(\mathbf{S})\mathbf{s} \rightarrow$ GFT coordinates $\tilde{\mathbf{x}} = \mathbf{diag}(\tilde{\mathbf{h}})\tilde{\mathbf{s}}$

Graph Fourier Transform (GFT)



- The Graph Fourier Transform (GFT) basis are the eigen-vectors of the GSO (generalization of the DFT case for the cycle graph).
- $\mathbf{S} = \mathbf{U}\mathbf{\Lambda}\mathbf{U}^{-1}$. The columns of $\mathbf{U}$ are the signals that are **invariant** w.r.t. the GSO $\mathbf{S}$

$$\tilde{\mathbf{x}} = \overbrace{\mathbf{U}^{-1}}^{GFT} \mathbf{x} \quad (3)$$

and the eigenvalues are the graph *frequencies* from low to high:

$$0 = \lambda_1 \leq \dots \leq \lambda_N.$$

- The **transfer function** of the GF is:

$$\tilde{\mathcal{H}}(\mathbf{\Lambda}) = \mathbf{diag}(\tilde{\mathbf{h}}) \quad \text{where} \quad [\tilde{\mathbf{h}}]_i = \tilde{h}_i = \sum_{\ell} h_{\ell} \lambda_i^{\ell}$$

node coordinates: $\mathbf{x} = \mathcal{H}(\mathbf{S})\mathbf{s} \rightarrow$ GFT coordinates $\tilde{\mathbf{x}} = \mathbf{diag}(\tilde{\mathbf{h}})\tilde{\mathbf{s}}$

Graph Fourier Transform (GFT)



- The Graph Fourier Transform (GFT) basis are the eigen-vectors of the GSO (generalization of the DFT case for the cycle graph).
- $\mathbf{S} = \mathbf{U}\mathbf{\Lambda}\mathbf{U}^{-1}$. The columns of $\mathbf{U}$ are the signals that are **invariant** w.r.t. the GSO $\mathbf{S}$

$$\tilde{\mathbf{x}} = \overbrace{\mathbf{U}^{-1}}^{GFT} \mathbf{x} \quad (3)$$

and the eigenvalues are the graph *frequencies* from low to high:

$$0 = \lambda_1 \leq \dots \leq \lambda_N.$$

- The **transfer function** of the GF is:

$$\tilde{\mathcal{H}}(\mathbf{\Lambda}) = \mathbf{diag}(\tilde{\mathbf{h}}) \quad \text{where} \quad [\tilde{\mathbf{h}}]_i = \tilde{h}_i = \sum_{\ell} h_{\ell} \lambda_i^{\ell}$$

node coordinates: $\mathbf{x} = \mathcal{H}(\mathbf{S})\mathbf{s} \rightarrow$ GFT coordinates $\tilde{\mathbf{x}} = \mathbf{diag}(\tilde{\mathbf{h}})\tilde{\mathbf{s}}$

Graph filters



- Thus, graph filters amplify or attenuate different graph frequencies exactly as classical filters do in time-domain DSP.

Transition

This viewpoint is the basis for modeling graph signals as *low-pass*, *bandlimited*, or approximately *low-rank* in the graph Fourier basis.

Outline



- 1 Review GSP
- 2 Unsupervised learning tasks
- 3 Low Rank Property of graph signals

Outline



- 1 Review GSP
- 2 Unsupervised learning tasks
- 3 Low Rank Property of graph signals**



Low rank data property

- Let $\mathbf{X} = (\mathbf{x}_1, \dots, \mathbf{x}_T)$. In a number of situations data indexed by graphs exhibit a sample covariance:

$$\hat{\mathbf{C}}_x = \frac{1}{T} \sum_{t=1}^T \mathbf{x}_t \mathbf{x}_t^\top = \frac{1}{T} \mathbf{X} \mathbf{X}^\top$$

that is **ill conditioned**.

- The GSP point of view is that $\mathbf{X} = (\mathbf{x}_1, \dots, \mathbf{x}_T)$ is potentially low rank because $\mathbf{U} \tilde{\mathbf{X}} \tilde{\mathbf{X}}^\top \mathbf{U}^\top = \mathbf{U} \text{diag}(\tilde{\mathbf{h}}) \tilde{\mathbf{C}}_s \text{diag}(\tilde{\mathbf{h}}) \mathbf{U}^\top$
- If the graph filter transfer function is sparse, irrespective of $\tilde{\mathbf{C}}_s$ only the non negligible values of $\tilde{h}(\lambda_i)$ will dominate the EVD of $\hat{\mathbf{C}}_x$



Low rank data property

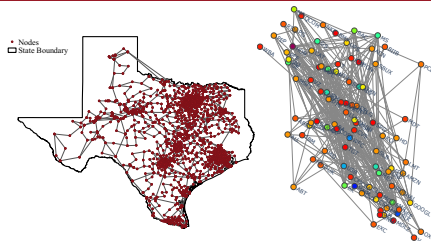
- Let $\mathbf{X} = (\mathbf{x}_1, \dots, \mathbf{x}_T)$. In a number of situations data indexed by graphs exhibit a sample covariance:

$$\hat{\mathbf{C}}_x = \frac{1}{T} \sum_{t=1}^T \mathbf{x}_t \mathbf{x}_t^\top = \frac{1}{T} \mathbf{X} \mathbf{X}^\top$$

that is **ill conditioned**.

- The GSP point of view is that $\mathbf{X} = (\mathbf{x}_1, \dots, \mathbf{x}_T)$ is potentially low rank because $\mathbf{U} \tilde{\mathbf{X}} \tilde{\mathbf{X}}^\top \mathbf{U}^\top = \mathbf{U} \text{diag}(\tilde{\mathbf{h}}) \tilde{\mathbf{C}}_s \text{diag}(\tilde{\mathbf{h}}) \mathbf{U}^\top$
- If the graph filter transfer function is sparse, irrespective of $\tilde{\mathbf{C}}_s$ only the non negligible values of $\tilde{h}(\lambda_i)$ will dominate the EVD of $\hat{\mathbf{C}}_x$

Examples



Social network

Power network

Financial network

Eigenvalues of the covariance matrix of Senate rollcall, voltage phasors and financial stock data. The low-pass nature is sufficient to create the low-rank property.

Orthogonal Projection



- You may recall from algebra the following result

Orthogonal projection

Let U_p be an orthonormal basis of a p -dimensional subspace of $\mathbb{R}^N$. Any vector in U_p can be expressed as $U_p v_p$, $v_p \in \mathbb{R}^p$.

The **orthogonal projection** $\hat{x}$ of x on the subspace U_p , is:

$$\hat{x} = U_p U_p^\top x.$$

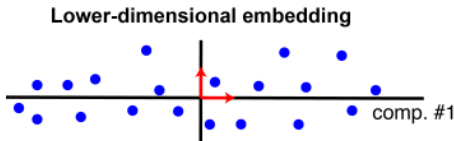
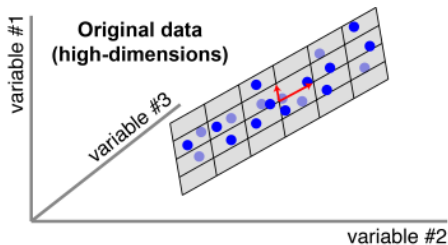
$\hat{x}$ is such that residual error $\hat{x} - x$ has minimum norm, i.e.

$$\hat{x} = \arg \min_{z: z=U_p v_p} \|z - x\|^2 \equiv U_p^\top x$$

Principal Component Analysis (PCA)



PCA tries to find the best subspace U_p that minimizes the residual. It is a generalization of linear regression: in the latter we try to find the straight line approximating $x(t) = at + b$, in PCA we have multi-dimensional data $\mathbf{X} = [\mathbf{x}_1, \dots, \mathbf{x}_T]$ and try to find an hyperplane



PCA solution



- Let us assume we removed the mean from the data $\mathbf{X}$ i.e.

$$\frac{1}{T} \sum_{t=1}^T \mathbf{x}_t = 0.$$
- As it turns out the solution is as follows:

Solution of batch PCA

The solution of

$$\min_{U_p} \sum_{t=1}^T \| (U_p U_p^\top - I) \mathbf{x}_t \|^2 \quad \text{subject to} \quad U_p^\top U_p = I$$

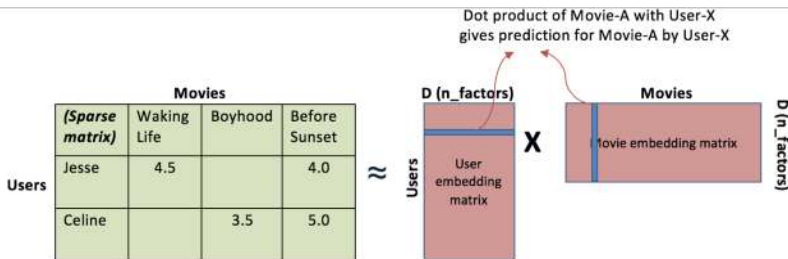
is the matrix of the p principal eigenvectors of $\frac{1}{T} \mathbf{X} \mathbf{X}^\top = \hat{\mathbf{C}}_x$,
 and the minimum cost is:

$$(U^\star)^H \hat{\mathbf{C}}_x U^\star = \text{diag}(\lambda_1, \dots, \lambda_p).$$

Modeling graph signals' low rank property



- PCA provides a method for dimensionality reduction and also for interpolation. For instance in collaborative filtering:

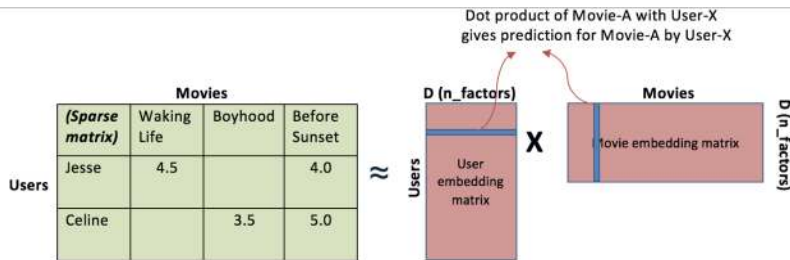


- When we model the data structure as $x(t) = \mathcal{H}(S)s$ the low rank structure must emerge from $\mathcal{H}(S)$. This is what we want to discuss next.

Modeling graph signals' low rank property



- PCA provides a method for dimensionality reduction and also for interpolation. For instance in collaborative filtering:



- When we model the data structure as $\mathbf{x}(t) = \mathcal{H}(\mathbf{S})\mathbf{s}$ the low rank structure must emerge from $\mathcal{H}(\mathbf{S})$. This is what we want to discuss next.

Modeling graph signals' low rank property



Recall that the GF

$$\mathcal{H}(\mathbf{S}) = \sum_{\ell=1}^L h_{\ell} \mathbf{S}^{\ell}$$

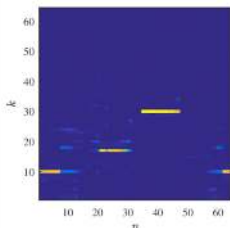
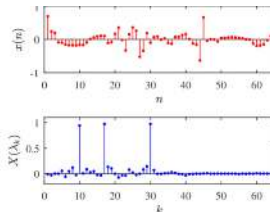
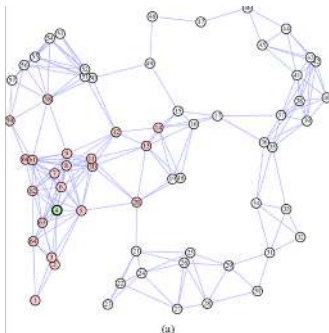
has a GFT response

$$\tilde{\mathbf{h}} : [\tilde{\mathbf{h}}]_i = \tilde{h}_i = \sum_{i=1}^N h_k \lambda_i^{\ell}$$

that shapes the GFT spectrum through the multiplicative property:

$$\mathbf{x} = \mathcal{H}(\mathbf{S}) \mathbf{s}$$

$$\Leftrightarrow \tilde{\mathbf{x}} = \text{diag}(\tilde{\mathbf{h}}) \tilde{\mathbf{s}}$$



Modeling graph signals' low rank property



- More specifically, where $|\tilde{h}_i| \ll 1$ it is expected that $|\tilde{x}_i| \ll 1$
- As a result, graph filter models can be classified as either **low-pass, band-pass, or high-pass**, or having a **sparse support**, depending on its graph frequency response, as we will see next.

GSP point of view

The vector found through PCA are, in fact, GFT components, i.e. eigenvectors of the GSO S .



Low-rank property

- Let us assume that we have independent and identically distributed *random experiments*:

$$\mathbf{x}_t = \mathcal{H}(\mathbf{S}) \mathbf{s}_t + \mathbf{w}_t \quad t = 1, \dots, T,$$

- Let $\mathbf{s}_\ell, \mathbf{w}_\ell$ are zero-mean white noise with $\mathbb{E}[\mathbf{s}_\ell \mathbf{s}_\ell^\top] = \mathbf{I}$, $\mathbb{E}[\mathbf{w}_\ell \mathbf{w}_\ell^\top] = \sigma^2 \mathbf{I}$. (This is not necessary but makes life easier)
- The first obvious property is that:

$$\hat{\mathbf{C}}_x \approx \mathbb{E}[\mathbf{x}_t \mathbf{x}_t^\top] = \mathcal{H}(\mathbf{S}) \mathcal{H}^\top(\mathbf{S}) + \sigma^2 \mathbf{I} \equiv \mathbf{U} \text{diag}(\tilde{\mathbf{h}})^2 \mathbf{U}^\top + \sigma^2 \mathbf{I}.$$

- **Low-pass graph signals** have the the first $\tilde{h}_1, \dots, \tilde{h}_k$ which dominate (corresponding to the small $\lambda_1, \dots, \lambda_k$)
- **High-pass graph signals** have $\tilde{h}_{N-k}, \dots, \tilde{h}_N$ which dominate
- **Sparse spectrum graph signals** have a certain subset $\mathcal{K}$ of the entries of $\tilde{\mathbf{h}}$ which dominate.



Low-rank property

- Let us assume that we have independent and identically distributed *random experiments*:

$$\mathbf{x}_t = \mathcal{H}(\mathbf{S}) \mathbf{s}_t + \mathbf{w}_t \quad t = 1, \dots, T,$$

- Let $\mathbf{s}_\ell$, $\mathbf{w}_\ell$ are zero-mean white noise with $\mathbb{E}[\mathbf{s}_\ell \mathbf{s}_\ell^\top] = \mathbf{I}$, $\mathbb{E}[\mathbf{w}_\ell \mathbf{w}_\ell^\top] = \sigma^2 \mathbf{I}$. (This is not necessary but makes life easier)
- The first obvious property is that:

$$\hat{\mathbf{C}}_x \approx \mathbb{E}[\mathbf{x}_t \mathbf{x}_t^\top] = \mathcal{H}(\mathbf{S}) \mathcal{H}^\top(\mathbf{S}) + \sigma^2 \mathbf{I} \equiv \mathbf{U} \text{diag}(\tilde{\mathbf{h}})^2 \mathbf{U}^\top + \sigma^2 \mathbf{I}.$$

- **Low-pass graph signals** have the the first $\tilde{h}_1, \dots, \tilde{h}_k$ which dominate (corresponding to the small $\lambda_1, \dots, \lambda_k$)
- **High-pass graph signals** have $\tilde{h}_{N-k}, \dots, \tilde{h}_N$ which dominate
- **Sparse spectrum graph signals** have a certain subset $\mathcal{K}$ of the entries of $\tilde{\mathbf{h}}$ which dominate.



Low-rank property

- Let us assume that we have independent and identically distributed *random experiments*:

$$\mathbf{x}_t = \mathcal{H}(\mathbf{S}) \mathbf{s}_t + \mathbf{w}_t \quad t = 1, \dots, T,$$

- Let $\mathbf{s}_\ell$, $\mathbf{w}_\ell$ are zero-mean white noise with $\mathbb{E}[\mathbf{s}_\ell \mathbf{s}_\ell^\top] = \mathbf{I}$, $\mathbb{E}[\mathbf{w}_\ell \mathbf{w}_\ell^\top] = \sigma^2 \mathbf{I}$. (This is not necessary but makes life easier)
- The first obvious property is that:

$$\hat{\mathbf{C}}_x \approx \mathbb{E}[\mathbf{x}_t \mathbf{x}_t^\top] = \mathcal{H}(\mathbf{S}) \mathcal{H}^\top(\mathbf{S}) + \sigma^2 \mathbf{I} \equiv \mathbf{U} \text{diag}(\tilde{\mathbf{h}})^2 \mathbf{U}^\top + \sigma^2 \mathbf{I}.$$

- **Low-pass graph signals** have the the first $\tilde{h}_1, \dots, \tilde{h}_k$ which dominate (corresponding to the small $\lambda_1, \dots, \lambda_k$)
- **High-pass graph signals** have $\tilde{h}_{N-k}, \dots, \tilde{h}_N$ which dominate
- **Sparse spectrum graph signals** have a certain subset $\mathcal{K}$ of the entries of $\tilde{\mathbf{h}}$ which dominate.

GFT and PCA



- Since PCA yields the dominant eigenvectors of $\hat{\mathbf{C}}_x$, then it is clear that those are the same as the Graph Frequencies $\mathbf{U}_{\mathcal{K}}$ corresponding to the dominant entries of $\tilde{\mathbf{h}}$
- If it is a low-pass signal, then the low GFT frequencies will dominate: they are the *least significant eigenvectors* of the GSO $\mathbf{S}$, i.e. they correspond to $0 = \lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_k$
 - Recall, the lowest eigenvectors of the GSP are those we use for spectral clustering!

GFT and PCA relation

Since the PCA basis is one to one with GFT components, if we know $\mathbf{S}$ we know what the principal components are and, if we do not know $\mathbf{S}$ we know how they relate to the graph.

GFT and PCA



- Since PCA yields the dominant eigenvectors of $\hat{\mathbf{C}}_x$, then it is clear that those are the same as the Graph Frequencies $\mathbf{U}_{\mathcal{K}}$ corresponding to the dominant entries of $\tilde{\mathbf{h}}$
- If it is a low-pass signal, then the low GFT frequencies will dominate: they are the *least significant eigenvectors* of the GSO $\mathbf{S}$, i.e. they correspond to $0 = \lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_k$
 - Recall, the lowest eigenvectors of the GSP are those we use for spectral clustering!

GFT and PCA relation

Since the PCA basis is one to one with GFT components, if we know $\mathbf{S}$ we know what the principal components are and, if we do not know $\mathbf{S}$ we know how they relate to the graph.

Outline



3 Low Rank Property of graph signals

- Sampling and interpolation theorem
- From graph data to graph topology

Sampling theory from GSP



- **Basic intuition:** Assume $\tilde{x}_t$ is dominated by the subset of components $\mathcal{K}$ that are the entries of $[\tilde{x}_t]_{\mathcal{K}}$. Then

$$\hat{x}_t = \mathbf{U}_{\mathcal{K}}[\tilde{x}_t]_{\mathcal{K}} \approx x_t,$$

where $\mathbf{U}_{\mathcal{K}}$ contains the the subset $\mathcal{K}$ of columns from the GFT basis and $[\tilde{x}_t]_{\mathcal{K}}$ contains the entries of $\tilde{x}_t$ in the set of indexes $\mathcal{K}$.

- In the case of lowpass graph signals $\mathcal{K} = \{1, \dots, k\}$
- The GSP sampling theorem simply states that, if the graph signal is a linear combinatio of relatively few GFT components, then we can decimate x_t over a subset of entries (samples) $\mathcal{M} \subset \mathcal{V}$ and it is possible to obtain good reconstruction accuracy as long as $|\mathcal{K}| \leq \mathcal{M}$.

Sampling theory from GSP



- **Basic intuition:** Assume $\tilde{x}_t$ is dominated by the subset of components $\mathcal{K}$ that are the entries of $[\tilde{x}_t]_{\mathcal{K}}$. Then

$$\hat{x}_t = \mathbf{U}_{\mathcal{K}}[\tilde{x}_t]_{\mathcal{K}} \approx x_t,$$

where $\mathbf{U}_{\mathcal{K}}$ contains the the subset $\mathcal{K}$ of columns from the GFT basis and $[\tilde{x}_t]_{\mathcal{K}}$ contains the entries of $\tilde{x}_t$ in the set of indexes $\mathcal{K}$.

- In the case of lowpass graph signals $\mathcal{K} = \{1, \dots, k\}$
- The GSP sampling theorem simply states that, if the graph signal is a linear combinatio of relatively few GFT components, then we can decimate x_t over a subset of entries (samples) $\mathcal{M} \subset \mathcal{V}$ and it is possible to obtain good reconstruction accuracy as long as $|\mathcal{K}| \leq \mathcal{M}$.

Sampling theory from GSP



- **Basic intuition:** Assume $\tilde{x}_t$ is dominated by the subset of components $\mathcal{K}$ that are the entries of $[\tilde{x}_t]_{\mathcal{K}}$. Then

$$\hat{x}_t = \mathbf{U}_{\mathcal{K}}[\tilde{x}_t]_{\mathcal{K}} \approx x_t,$$

where $\mathbf{U}_{\mathcal{K}}$ contains the the subset $\mathcal{K}$ of columns from the GFT basis and $[\tilde{x}_t]_{\mathcal{K}}$ contains the entries of $\tilde{x}_t$ in the set of indexes $\mathcal{K}$.

- In the case of lowpass graph signals $\mathcal{K} = \{1, \dots, k\}$
- The GSP sampling theorem simply states that, if the graph signal is a linear combinatio of relatively few GFT components, then we can decimate x_t over a subset of entries (samples) $\mathcal{M} \subset \mathcal{V}$ and it is possible to obtain good reconstruction accuracy as long as $|\mathcal{K}| \leq \mathcal{M}$.

Sampling theory from GSP



Suppose $\mathbf{x}_t \approx \hat{\mathbf{x}}_t = \mathbf{U}_{\mathcal{K}}[\tilde{\mathbf{x}}_t]_{\mathcal{K}} = \mathbf{U}_{\mathcal{K}}\mathbf{U}_{\mathcal{K}}^{\top}\mathbf{x}_t$. Note that the GFT of $\hat{\mathbf{x}}$ is zero for all graph frequencies $i \notin \mathcal{K}$.

Vertex-limiting (sampling) operator

Measurements from nodes $i \in \mathcal{M}$ $[\mathbf{x}_t]_{\mathcal{M}} = \mathbf{P}_{\mathcal{M}}^{\top}\mathbf{x}_t$
 Vertex-limiting operator. $\mathcal{D}_{\mathcal{M}} = \mathbf{P}_{\mathcal{M}}\mathbf{P}_{\mathcal{M}}^{\top}$

- Note: $[\mathbf{x}_t]_{\mathcal{M}} = \mathbf{P}_{\mathcal{M}}^{\top}\mathbf{x}_t \approx \mathbf{P}_{\mathcal{M}}^{\top}\mathbf{U}_{\mathcal{K}}[\tilde{\mathbf{x}}_t]_{\mathcal{K}} \rightarrow [\tilde{\mathbf{x}}_t]_{\mathcal{K}} = (\mathbf{P}_{\mathcal{M}}^{\top}\mathbf{U}_{\mathcal{K}})^{\dagger}[\mathbf{x}_t]_{\mathcal{M}}$
- The **reconstruction formula** is:

$$\hat{\mathbf{x}}_t = \mathbf{U}_{\mathcal{K}}(\mathbf{P}_{\mathcal{M}}^{\top}\mathbf{U}_{\mathcal{K}})^{\dagger}[\mathbf{x}_t]_{\mathcal{M}} \quad (4)$$

$|\mathcal{M}| \geq |\mathcal{K}|$ is a necessary condition and the optimal sampling pattern choose rows of $\mathbf{U}_{\mathcal{K}}$ so as to maximize the **smallest singular value**, $\sigma_{\min}(\mathcal{D}_{\mathcal{M}}\mathbf{U}_{\mathcal{K}})$.

Interpolation w.o. GSP: Matrix completion



- Let $\mathbf{X}^{\text{samp}}$ be a matrix that contains the set of entries of the vectors $\mathbf{x}_t$ that are known $\{[\mathbf{x}_1]_{\mathcal{M}_1}, \dots, [\mathbf{x}_T]_{\mathcal{M}_T}\}$ and zeros when they are missing.
- Assume we know tend to be well approximated in a subspace of dimension $p > n$, even though it is not exactly p -dimensional

Matrix Completion Problem

Let $\|\mathbf{X}\|^\star$ denote the **nuclear norm**, i.e. the sum of the magnitudes of the singular values of $\mathbf{X}$

$$\min_{\mathbf{X} \in \mathbb{R}^{n \times m}} \|\mathcal{M}(\mathbf{X}) - \mathbf{X}^{\text{samp}}\|_F^2 + \gamma \left\{ \|\mathbf{X}\|^\star + \sum_{t=2}^m \|\mathbf{x}_t - \mathbf{x}_{t-1}\|_2^2 \right\}.$$

- The regularization term $\{\|\mathbf{X}\|^\star + \sum_{t=2}^m \|\mathbf{x}_t - \mathbf{x}_{t-1}\|_2^2\}$ tends to promote having just a few dominant singular values

Smoothness and interpolation with GSP

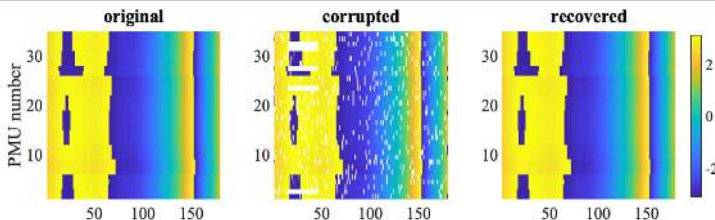


Cornell University

- For **Graph-Temporal processes** that are Low-pass using the GSO yields better results,:

$$\min_{\mathbf{X} \in \mathbb{R}^{n \times m}} \|\mathcal{M}(\mathbf{X}) - \mathbf{X}^{\text{samp}}\|_F^2 + \gamma \left\{ \sum_{t=1}^T \mathbf{x}_t^\top \mathbf{S} \mathbf{x}_t + \sum_{t=2}^m \|\mathbf{x}_t - \mathbf{x}_{t-1}\|_2^2 \right\}.$$

- The regularization $\sum_{t=1}^T \mathbf{x}_t^\top \mathbf{S} \mathbf{x}_t$ seeks to reduce spatial variations main benefit of this approach is that it can recover **entirely missing rows** of data, unlike matrix completion



Outline



3 Low Rank Property of graph signals

- Sampling and interpolation theorem
- From graph data to graph topology



Smoothness and Graph Learning

- For low-pass graph signals with independent samples, in general we can formulate the following problem. **Graph-Topology** learning formulation:

$$\begin{aligned} \min_{\mathbf{z}_t, t=1, \dots, T, \mathbf{S}} \quad & \frac{1}{m} \sum_{t=1}^T \left(\frac{1}{\sigma^2} \|\mathbf{z}_t - \mathbf{x}_t\|_2^2 + \mu \mathbf{z}_t^\top \mathbf{S} \mathbf{z}_t \right) \\ \text{s.t.} \quad & S_{ji} = S_{ij} \leq 0, \forall i \neq j, \mathbf{S} \mathbf{1} = \mathbf{0}, \end{aligned}$$

finding the Laplacian that best **smoothens** the signal ($\mathbf{z}_t$ is a smooth signal close to the observation $\mathbf{x}_t$ which may be noisy).

- Note that if the previous interpolation problem is solved w.r.t. both $\mathbf{X}$ and $\mathbf{S}$ we obtain a similar formulation, where also temporal smoothness is accounted for

$$\begin{aligned} \min_{\mathbf{Z} \in \mathbb{R}^{n \times m}, \mathbf{S}} \quad & \|\mathcal{M}(\mathbf{Z}) - \mathbf{X}^{\text{samp}}\|_F^2 + \gamma \left(\sum_{t=1}^T \mathbf{z}_t^\top \mathbf{S} \mathbf{z}_t + \sum_{t=2}^m \|\mathbf{z}_t - \mathbf{z}_{t-1}\|_2^2 \right) \\ \text{s.t.} \quad & S_{ji} = S_{ij} \leq 0, \forall i \neq j, \mathbf{S} \mathbf{1} = \mathbf{0}, \end{aligned}$$



Smoothness and Graph Learning

- For low-pass graph signals with independent samples, in general we can formulate the following problem. **Graph-Topology** learning formulation:

$$\begin{aligned} \min_{\mathbf{z}_t, t=1, \dots, T, \mathbf{S}} \quad & \frac{1}{m} \sum_{t=1}^T \left(\frac{1}{\sigma^2} \|\mathbf{z}_t - \mathbf{x}_t\|_2^2 + \mu \mathbf{z}_t^\top \mathbf{S} \mathbf{z}_t \right) \\ \text{s.t.} \quad & S_{ji} = S_{ij} \leq 0, \forall i \neq j, \mathbf{S} \mathbf{1} = \mathbf{0}, \end{aligned}$$

finding the Laplacian that best **smoothens** the signal ($\mathbf{z}_t$ is a smooth signal close to the observation $\mathbf{x}_t$ which may be noisy).

- Note that if the previous interpolation problem is solved w.r.t. both $\mathbf{X}$ and $\mathbf{S}$ we obtain a similar formulation, where also temporal smoothness is accounted for

$$\begin{aligned} \min_{\mathbf{Z} \in \mathbb{R}^{n \times m}, \mathbf{S}} \quad & \|\mathcal{M}(\mathbf{Z}) - \mathbf{X}^{\text{samp}}\|_F^2 + \gamma \left(\sum_{t=1}^T \mathbf{z}_t^\top \mathbf{S} \mathbf{z}_t + \sum_{t=2}^m \|\mathbf{z}_t - \mathbf{z}_{t-1}\|_2^2 \right) \\ \text{s.t.} \quad & S_{ji} = S_{ij} \leq 0, \forall i \neq j, \mathbf{S} \mathbf{1} = \mathbf{0}, \end{aligned}$$

Blind community detection (CD)



- We know that the principal components of C_y will be dominated by a rank- k component spanned by U_k which are the **least significant eigenvectors** of the Laplacian

Blind CD algorithm

Upon observing m low-pass graph signals:

- (i) find the top- k $\hat{U}_k \in \mathbb{R}^{n \times k}$ of $\hat{C}_y = \frac{1}{m} \sum_{\ell=1}^m \mathbf{y}_\ell \mathbf{y}_\ell^\top$;
- (ii) apply k -means on the rows of $\hat{U}_k$.

- If we denote the detected communities as $\hat{\mathcal{N}}_1, \dots, \hat{\mathcal{N}}_k$, then

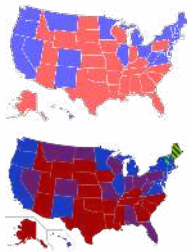
$$F(\hat{\mathcal{N}}_1, \dots, \hat{\mathcal{N}}_k; U_k) - F^\star = \mathcal{O}(\eta_k + \sigma + m^{-1/2}).$$

- In other words, the BlindCD approaches the performance of SC if the graph filter is k -low-pass with $\eta_k \ll 1$, the observation noise σ and the observation noise σ is small.

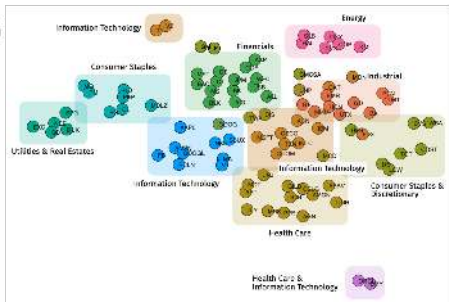
Blind CD in Opinion Dynamics



(a)



(b)



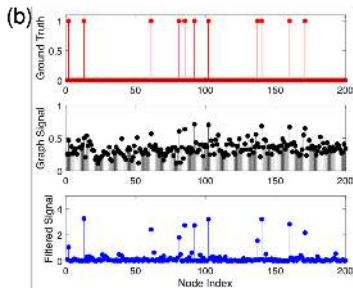
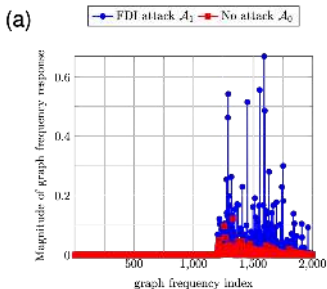
- (a) US Senate from the 115th US Congress (2017-2019)
 $m = 502$ voting rounds - grouping senators from each state.
- (b) Daily return data of stocks in the S&P100 index from Feb. 2013.

Anomaly detection with Low-pass GSP



- We can cast the anomaly detection as a canonical hypothesis test:
 - **Null Hypothesis $\mathcal{A}_0$** : The high frequency components of the graph signal are small
 - **Positive Hypothesis $\mathcal{A}_1$** : The high frequency graph signal components are large due to a sudden change or inconsistent data

$$x_\ell = \begin{cases} \mathcal{H}(\mathbf{S})s_\ell, & \text{under } \mathcal{A}_0, \\ \mathcal{H}(\mathbf{S})s_\ell + w_\ell & \text{under } \mathcal{A}_1, \end{cases} \Rightarrow \Gamma_\ell = \|\mathcal{H}_{\text{HPF}}(\mathbf{S})x_\ell\| \underset{\mathcal{A}_1}{\overset{\mathcal{A}_0}{\leq}} \delta.$$



Conclusion



- We introduced a class of unsupervised algorithms that used data to make inferences and are based on the GSP framework.
 - Least square reconstruction of graph signals from sparse nodal measurements and optimum sampling design
 - Matrix completion versus regularized reconstruction of low-pass graph temporal signals
 - Joint network inference and data reconstruction
 - Community detection from graph signals
 - Anomaly detection
- Next we discuss how to perform supervised inference with Graph Convolutional Neural Networks