

Mapping Flows on Complex Networks with Graph Neural Networks

NetSciX 2023, Venice, Italy

Christopher Blöcker^{1,2}, Chester Tan¹, Ingo Scholtes^{1,2}

christopher.bloecker@uzh.ch

¹Chair of Machine Learning for Complex Networks
Center for Artificial Intelligence and Data Science
Julius-Maximilians-Universität Würzburg
Würzburg, Germany



²Data Analytics Group
Department of Informatics
University of Zurich
Zurich, Switzerland



**Universität
Zürich**^{UZH}

too long; didn't listen

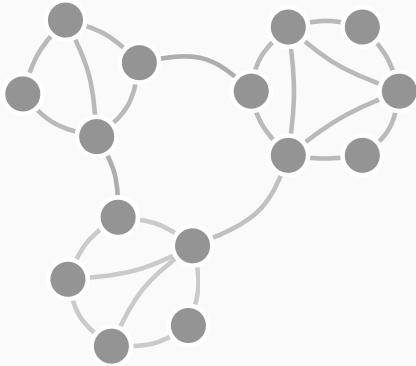
map equation + graph neural networks

- Community-detection methods typically rely on optimising objective functions without leveraging recent advances in deep learning
- We express the map equation in differentiable tensor form for optimisation with different (G)NN architectures and gradient descent on GPUs
- It works well for community detection in synthetic and real networks, naturally detects overlapping communities, and can incorporate node features

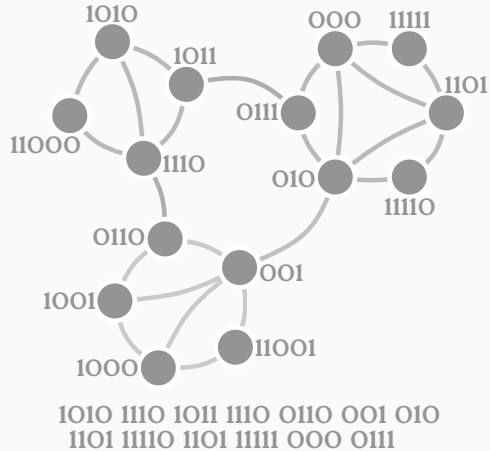
The Map Equation

→ Rosvall and Bergstrom; PNAS 2008

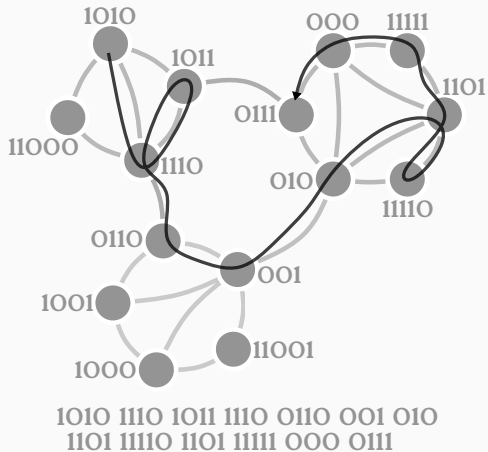
The Map Equation



The Map Equation

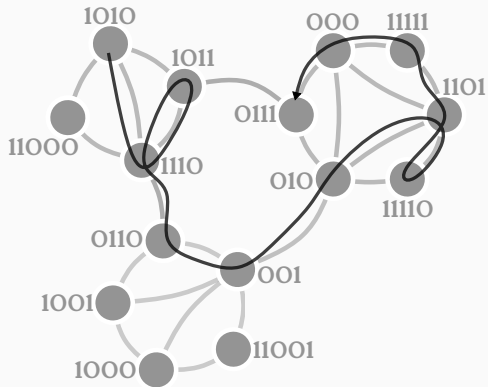


The Map Equation



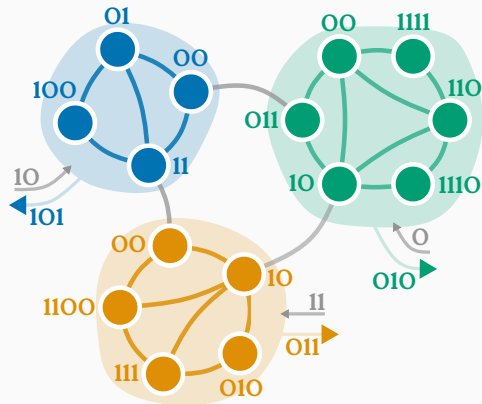
$$L(M) = H(P) = - \sum_p p \log_2 p \text{ [bits]}$$

The Map Equation

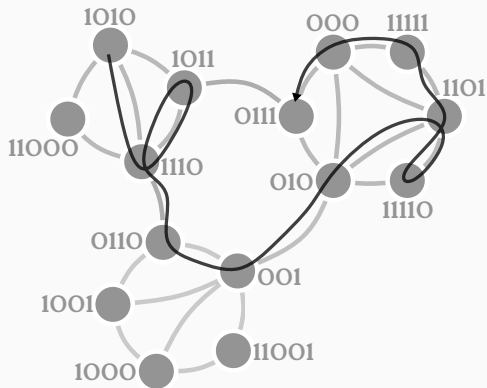


1010 1110 1011 1110 0110 001 010
1101 1110 1101 1111 000 0111

$$L(M) = H(P) = - \sum_p p \log_2 p \text{ [bits]}$$

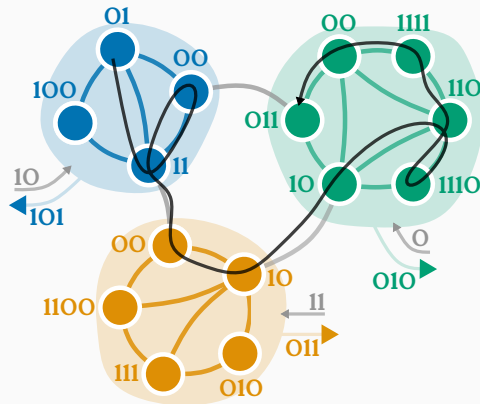


The Map Equation



1010 1110 1011 1110 0110 001 010
1101 1110 1101 1111 000 011

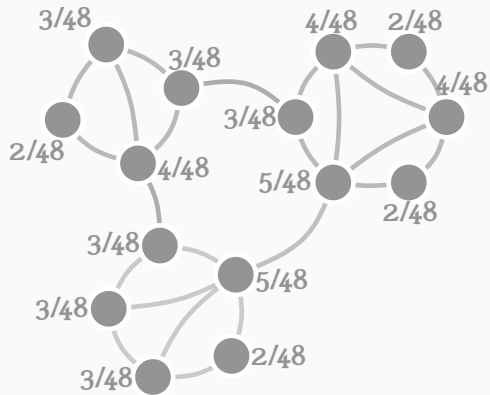
$$L(M) = H(P) = - \sum_p p \log_2 p \text{ [bits]}$$



1001 11 00 11 1011 00 10 0110 10
110 1110 110 1111 00 011

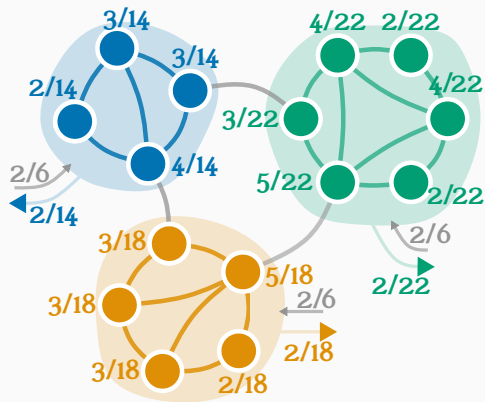
$$L(M) = qH(Q) + \sum_{m \in M} p_m H(P_m) \text{ [bits]}$$

The Map Equation



3.84 bits

$$L(M) = H(P) = - \sum_p p \log_2 p \text{ [bits]}$$



3.05 bits

$$L(M) = qH(Q) + \sum_{m \in M} p_m H(P_m) \text{ [bits]}$$

The Map Equation Goes Neural

The Map Equation Goes Neural

- We introduce a soft cluster assignment matrix, making node assignments partial
 - differentiable map equation
 - overlapping modules
- We implement the differentiable map equation in tensor form
 - loss term for (G)NNs
- Optimisation with gradient descent means differentiating the codelength with respect to soft cluster assignments
 - update all assignments a bit instead of changing one binary assignment
- No regularisation beyond using the map equation

$$S = \text{softmax}(\text{GNN}(A, X)) \quad \underset{S}{\text{argmin}} L(S, F)$$

Going Neural

The map equation can be rewritten as

→ Rosvall et al.; Eur. Phys. J. 2009

$$\left(\sum_{m \in M} q_m\right) \log_2 \left(\sum_{m \in M} q_m\right) - 2 \sum_{m \in M} q_m \log_2 q_m - \sum_{u \in V} p_u \log_2 p_u + \sum_{m \in M} \left(m_{\text{exit}} + \sum_{u \in m} p_u\right) \log_2 \left(m_{\text{exit}} + \sum_{u \in m} p_u\right)$$

We implement the four parts

$$e_1 = \sum_{a,b} E_{ab} - \sum_{a,b} \delta_{ab} E_{ab} \quad (e_2)_a = \sum_b (E_{ab} - \delta_{ab} E_{ab}) \quad (e_3)_i = p_i \quad (e_4)_b = \sum_a (E_{ab} - \delta_{ab} E_{ab}) + \sum_a E_{ab}$$

where

- S is the soft cluster assignment matrix,
- p is the vector of node visit rates.
- F is the flow matrix,
- $E_{ab} = \sum_{i,j} S_{ia}^T F_{ij} S_{jb}$ summarises the flow from module a to b ,

And put them back together,

$$L(S, F) = e_1 \log_2 e_1 - 2 \sum_{a,b} ((e_2)_a \delta_{ab} \log_2 (e_2)_b) - \sum_{i,j} ((e_3)_i \delta_{ij} \log_2 (e_3)_j) + \sum_{a,b} ((e_4)_a \delta_{ab} \log_2 (e_4)_b)$$

Evaluation

Community Detection

We use Infomap and different (G)NN architectures

- Simple linear layer and MLP
- GCN_t with $t \in \{1, 2\}$ layers $\rightarrow$ Kipf and Welling; arXiv 2016
- GIN_t with $t \in \{1, 2, 3\}$ layers $\rightarrow$ Xu et al.; arXiv 2018
- GNN_t with $t \in \{1, 2\}$ layers: message passing + skip-connections between MLP layers
- learning rate $\{10^{-1}, 10^{-2}, \dots, 10^{-6}\}$

Datasets

25 LFR networks with planted (hard) communities for each combination of

- $n \in \{100, 1000\}$ nodes, $X = I$
- average degree $k \in \{\ln n, 2 \ln n\}$
- maximum degree $k_{\max} = 2\sqrt{n}$
- mixing $\mu \in \{0.1, 0.2, \dots, 0.8\}$

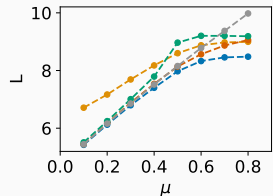
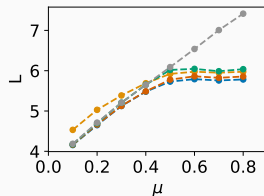
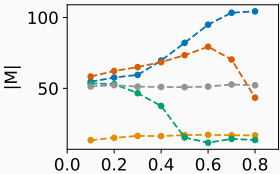
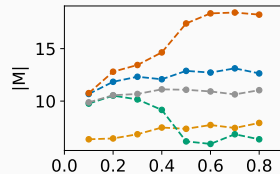
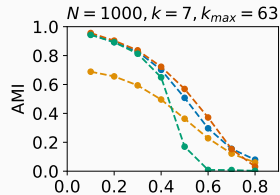
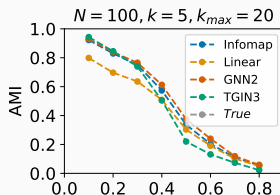
Results – undirected

We measure

- AMI against ground truth
- Number of modules $|M|$
- Codelength L

Main points

- Similar performance to Infomap, slightly outperforming it (AMI)
- performance $\sim$ (G)NN expressivity
→ over/underfitting ($|M|$)
- Infomap achieves lowest L
- consider all three measures!
→ similar AMI but different $|M|$
→ Smaller L always better?



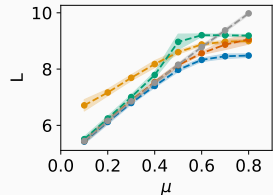
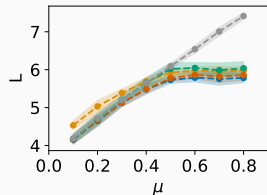
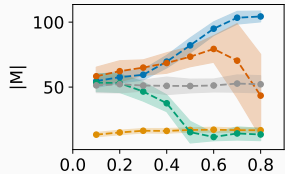
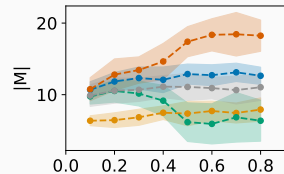
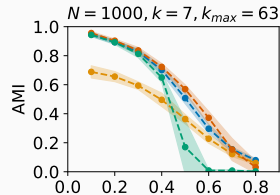
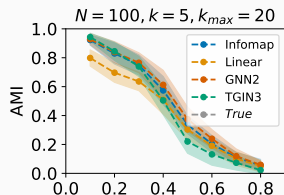
Results – undirected

We measure

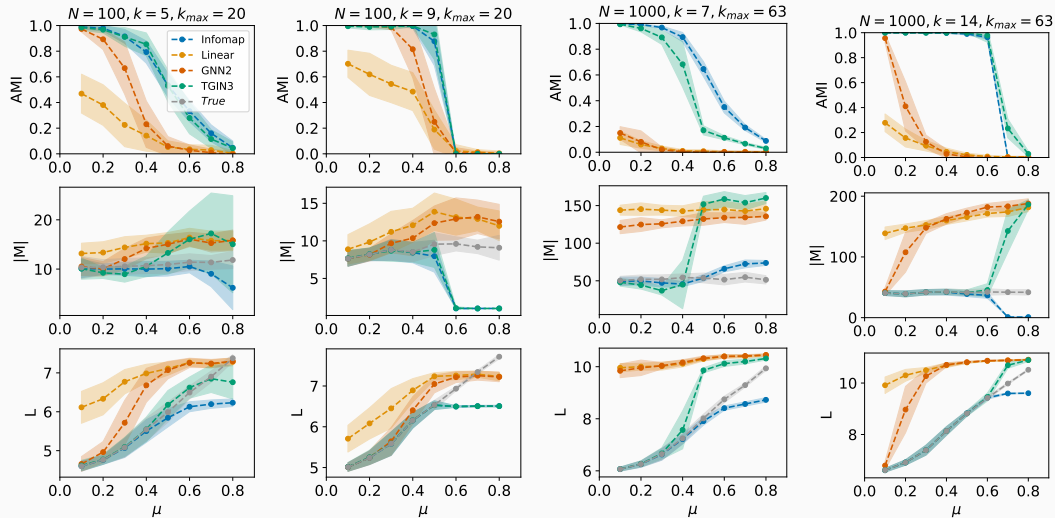
- AMI against ground truth
- Number of modules $|M|$
- Codelength L

Main points

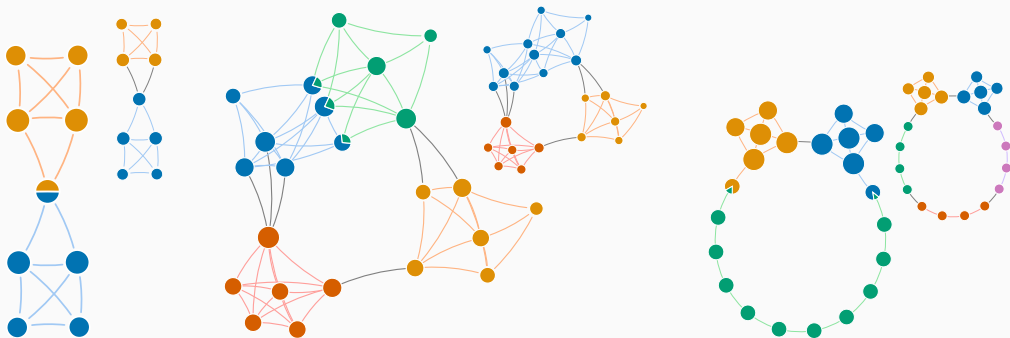
- Similar performance to Infomap, slightly outperforming it (AMI)
- performance $\sim$ (G)NN expressivity
→ over/underfitting ($|M|$)
- Infomap achieves lowest L
- consider all three measures!
→ similar AMI but different $|M|$
→ Smaller L always better?



Results – directed



Results



Conclusion

Conclusion

Motivation

Leverage advances from deep learning for community detection (GPUs!)

Our Approach

We make the map equation differentiable by introducing soft cluster assignments, and optimise it with (G)NNs and gradient descent.

Results

- Community-detection performance, overfitting, and underfitting depend on the chosen (G)NN architecture
- We naturally detect overlapping communities
- Benefit: enables fast experimentation with map equation adoptions!

Thank you for your attention!
Check out our preprint: [arXiv:2310.01144](https://arxiv.org/abs/2310.01144)



Universität
Zürich^{UZH}