

Directed Acyclic Graph (DAG)-Informed Structure Learning from Multi-Dimensional Point Processes

Chunming Zhang¹ (czhang3@wisc.edu) Muhong Gao² Shengji Jia³

¹University of Wisconsin-Madison



²Chinese Academy of Sciences

³Shanghai Lixin University of Accounting and Finance

1. Introduction: Multivariate Temporal Point Process

■ Multivariate Temporal Point Process:

- Random occurrences of a specific type of event (e.g., neuron spike firing, contagious disease incidence, e-commerce transactions, or social media postings) over time, represented as $\{T_1, \dots, T_d\}$, recorded at d nodes.
- At each node j in the node set $\mathcal{V} = \{1, \dots, d\}$,
 $T_j = (T_{j,1}, \dots, T_{j,N_j})$ with $0 < T_{j,1} < \dots < T_{j,N_j} \leq T$, (1)
 signifies a sequence of time points $T_{j,\ell}$, representing the occurrence time of the ℓ th event during an experiment of duration T .

■ **Key objective:** Learn the dependency structure among nodes and represent it as a **DAG** from the d sequences of event times.

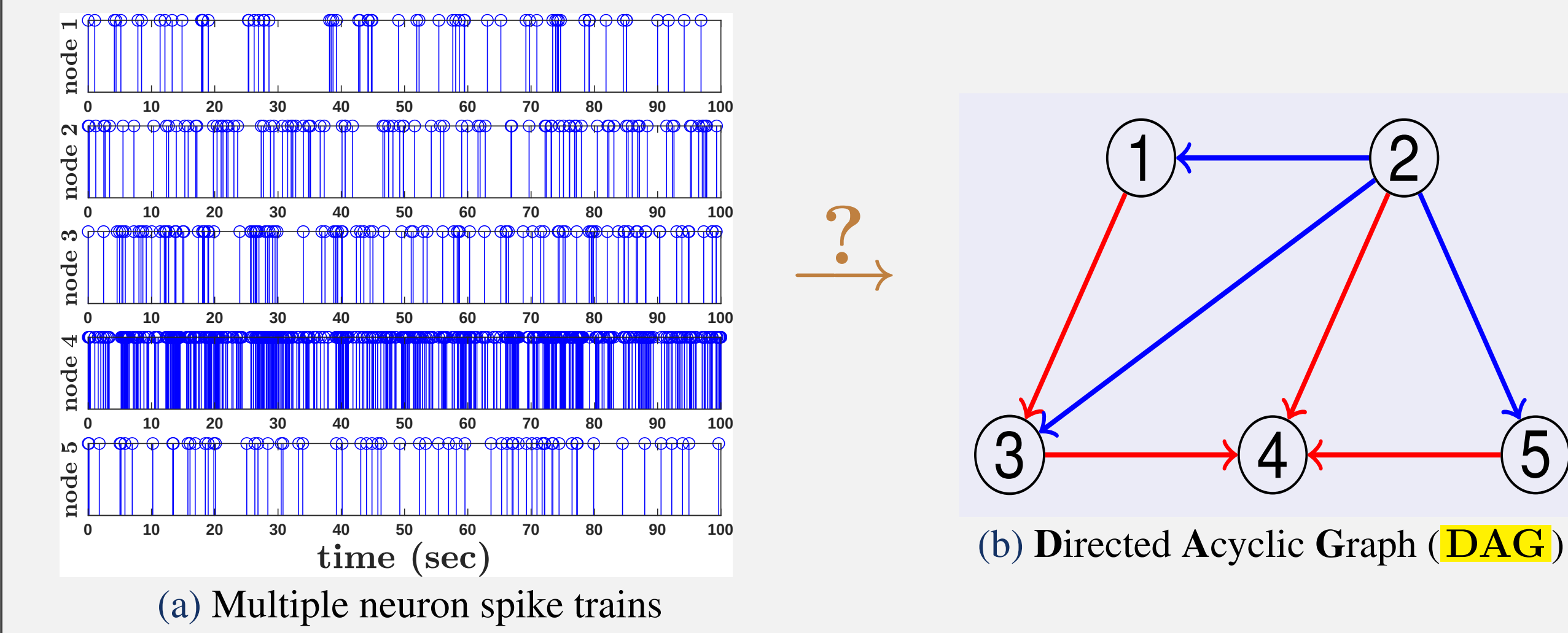


Figure: [Multiple Neuron Spike Train Data] Each node of the network graph in panel (b) corresponds to a point process in panel (a). Arrows indicate interactions (red: excitatory; blue: inhibitory).

2. Generative Modeling of CIFs for Causal Graphs

■ Our continuous-time model for the conditional intensity function:

$$\lambda_j(t | \mathcal{F}_t) = \exp(w_{0,j} + \sum_{i=1}^d w_{i,j} x_i(t)) \quad (2)$$

■ Negative log-likelihood function for point process data T_j :

$$\mathcal{L}_{j,T}(\mathbf{w}_{\cdot,j}) = -T^{-1} \int_0^T [\mathbf{w}_{\cdot,j}^\top \mathbf{x}(t-) dN_j(t) - \exp\{\mathbf{w}_{\cdot,j}^\top \mathbf{x}(t)\} dt]. \quad (3)$$

■ Proposed ‘DAG-Constrained Regularized Estimator’:

$$\mathbf{W} = \arg \min_{\mathbf{W} \in \mathbb{R}^{(d+1) \times d}} \mathcal{L}(\mathbf{W}) + \mathcal{P}(\mathbf{W}; \boldsymbol{\eta}), \quad (4)$$

subject to $\mathcal{G}(\mathbf{W})$ being a DAG,

where $\mathbf{W} = (w_{i,j}) \in \mathbb{R}^{d \times d}$ is the weighted adjacency matrix,

$$\mathcal{L}(\mathbf{W}) = \sum_{j=1}^d \mathcal{L}_{j,T}(\mathbf{w}_{\cdot,j}), \quad (5)$$

$$\mathcal{P}(\mathbf{W}; \boldsymbol{\eta}) = \sum_{1 \leq i \neq j \leq d} \eta_{i,j} |w_{i,j}|, \quad \text{with } \eta_{i,j} \geq 0. \quad (6)$$

■ Question:

- How can we solve (4) under the combinatorial DAG-constraint?
- What are the statistical properties of the solution?

3. Optimization Algorithm for DAG-Constrained Estimator

■ Characterizing a DAG (Zheng et al., 2018, NeurIPS):

$$\mathcal{G}(\mathbf{W}) \text{ is a DAG} \iff h(\mathbf{W}) = 0,$$

where

$$h(\mathbf{W}) := \text{trace}\{\exp(\mathbf{W} \circ \mathbf{W})\} - d. \quad (7)$$

■ Our (4) becomes a continuous equality-constrained program:

$$\min_{\mathbf{W} \in \mathbb{R}^{(d+1) \times d}} \mathcal{L}(\mathbf{W}) + \mathcal{P}(\mathbf{W}; \boldsymbol{\eta}), \quad (8)$$

subject to $h(\mathbf{W}) = 0$.

■ The augmented Lagrangian (AL) function for problem (8):

$$L(\mathbf{W}, \alpha; \rho) = \mathcal{L}(\mathbf{W}) + \mathcal{P}(\mathbf{W}; \boldsymbol{\eta}) + \alpha h(\mathbf{W}) + 2^{-1} \rho h^2(\mathbf{W}), \quad (9)$$

where $\alpha \in \mathbb{R}$ is the Lagrange multiplier, and $\rho > 0$ is the augmentation parameter.

■ The classical AL method is a dual ascent algorithm: For a fixed ρ ,

- Given α , solve

$$\bar{\mathbf{W}}_\alpha = \arg \min_{\mathbf{W} \in \mathbb{R}^{(d+1) \times d}} L(\mathbf{W}, \alpha; \rho). \quad (10)$$

- Define the dual function as the Lagrange multiplier α ,

$$D(\alpha) = L(\bar{\mathbf{W}}_\alpha, \alpha; \rho), \quad (11)$$

and solve

$$\bar{\alpha} = \arg \max_{\alpha \in \mathbb{R}} D(\alpha). \quad (12)$$

■ The classical AL algorithm for updating $\mathbf{W}$, α , and ρ :

- Update $\mathbf{W}$** , by solving the unconstrained subproblem:

$$\mathbf{W}^{(k+1)} = \arg \min_{\mathbf{W} \in \mathbb{R}^{(d+1) \times d}} L(\mathbf{W}, \alpha^{(k)}; \rho^{(k)}), \quad (13)$$

using the previous solution $\mathbf{W}^{(k)}$ as the initial value.

- Update α and ρ** according to the rules:

$$\alpha^{(k+1)} = \alpha^{(k)} + \beta_\alpha h(\mathbf{W}^{(k+1)}), \quad (14)$$

$$\rho^{(k+1)} = \beta_\rho \rho^{(k)}, \quad (15)$$

for appropriately chosen step sizes $\beta_\alpha > 0$ and $\beta_\rho \geq 1$.

■ Limitation of classical AL method: Slow convergence and high computational cost.

- Computing $h(\mathbf{W})$ in (14) is particularly expensive.
- Solving (13) is computationally intensive.

■ Proposed Flex-AL algorithm with flexible updates of α and ρ :

$$\max\{\alpha^{(k)}, \rho^{(k)}\} \rightarrow \infty \text{ as } k \rightarrow \infty, \quad (16)$$

and $\inf_{k \geq 0} \alpha^{(k)} > -\infty, \quad \inf_{k \geq 0} \rho^{(k)} > 0$.

■ **Example update:** $\alpha^{(k+1)} = \gamma_\alpha \alpha^{(k)}$ and $\rho^{(k+1)} = \gamma_\rho \rho^{(k)}$, where $\gamma_\alpha > 1$ and $\gamma_\rho > 1$, starting from $\alpha^{(0)} > 0$ and $\rho^{(0)} > 0$.

Global convergence of Flex-AL algorithm for solving (8):

Assume that the sequence of Lagrange multipliers $\{\alpha^{(k)}\}_{k \geq 0}$ and the sequence of augmentation parameters $\{\rho^{(k)}\}_{k \geq 0}$ satisfy condition (16). Then:

- Any limit point of the sequence of global minimizers $\{\mathbf{W}^{(k)}\}_{k \geq 1}$ in (13) is a global minimizer of (8).
- Moreover, assume condition A6' (the minimum eigenvalue of $\nabla^2 \mathcal{L}_{j,T}(\mathbf{w}_{\cdot,j})$ is bounded away from 0). Then $\{\mathbf{W}^{(k)}\}_{k \geq 1}$ has at least one limit point. Also, if (8) has a unique global minimizer $\bar{\mathbf{W}}$, then $\mathbf{W}^{(k)} \rightarrow \bar{\mathbf{W}}$ as $k \rightarrow \infty$.

4. Statistical Properties of Structure Learning Methods

■ Expanded estimators:

$$\bar{\mathbf{W}}_{\boldsymbol{\kappa}, \boldsymbol{\eta}} = \arg \min_{\mathbf{W} \in \mathbb{R}^{(d+1) \times d}; h(\mathbf{W}) \leq \boldsymbol{\kappa}} \mathcal{L}(\mathbf{W}) + \mathcal{P}(\mathbf{W}; \boldsymbol{\eta}). \quad (17)$$

- Scenario (a): Unregularized Estimation ($\boldsymbol{\eta} = 0$ in (17))**

– Case 1: non-DAG-constrained ($\boldsymbol{\kappa} = \infty$), unregularized ($\boldsymbol{\eta} = 0$):

– Case 2: DAG-constrained ($\boldsymbol{\kappa} = 0$), unregularized ($\boldsymbol{\eta} = 0$):

- Scenario (b): Regularized Estimation ($\boldsymbol{\eta} \neq 0$ in (17))**

– Case 3: non-DAG-constrained ($\boldsymbol{\kappa} = \infty$), regularized ($\boldsymbol{\eta} \neq 0$):

– Case 4: DAG-constrained ($\boldsymbol{\kappa} = 0$), regularized ($\boldsymbol{\eta} \neq 0$):

■ Takeaways:

- Without sparsity-regularization, incorporating the DAG constraint is necessary for DAG-recovery consistency.
- With an appropriate sparsity-regularization penalty, the DAG-constrained estimator achieves the same parameter-estimation and DAG-recovery consistencies as the non-DAG-constrained estimator, while improving empirical performance.

5. Numerical Experiments

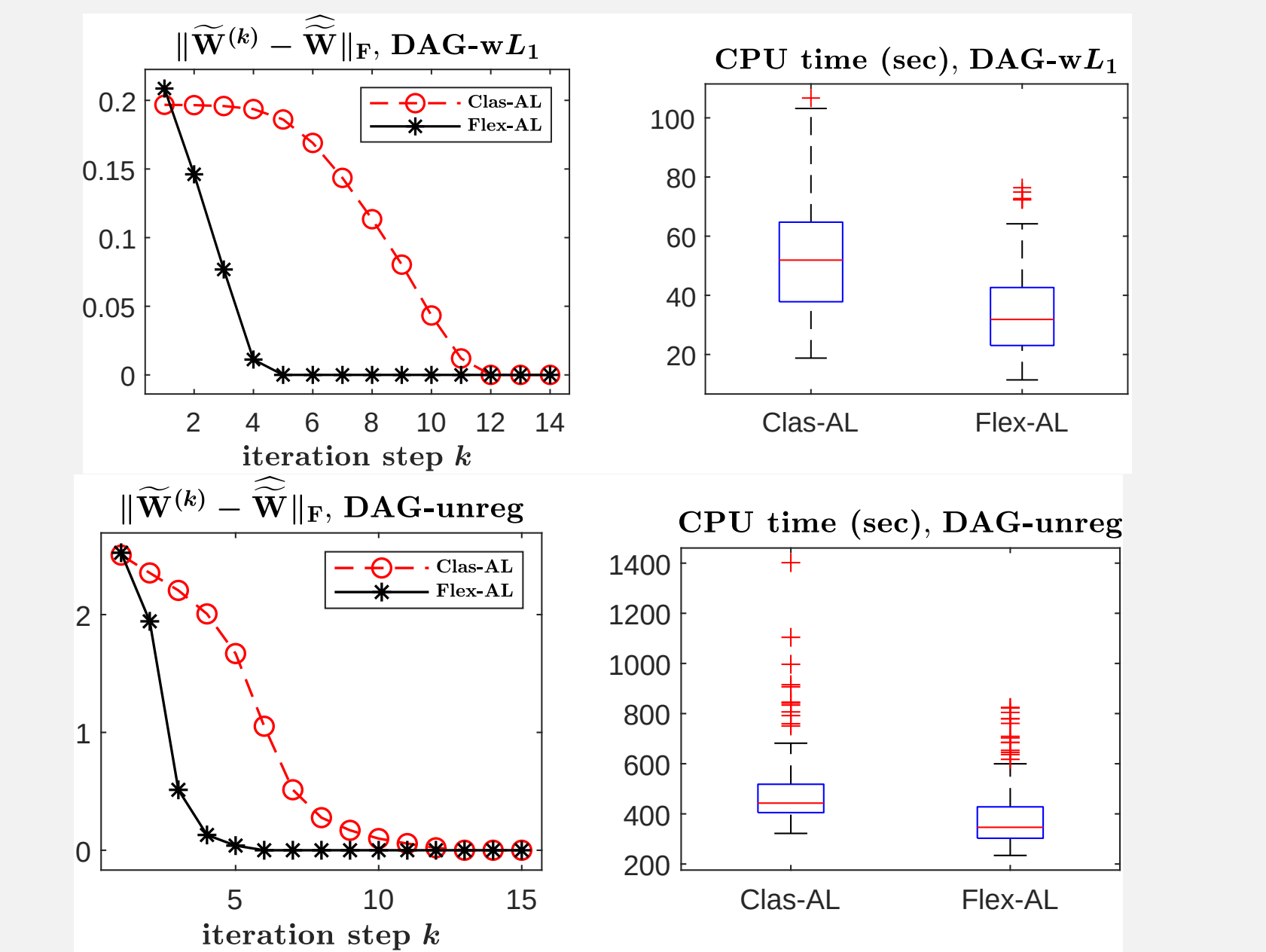


Figure: (Simulation Study: Network 3) Left panel: average approximation error $\|\mathbf{W}^{(k)} - \bar{\mathbf{W}}\|_F$ over 100 replicates versus the iteration step k ; right panel: boxplots of runtime for estimators $\mathbf{W}$ across replications. Top row: ‘DAG-wL1’; bottom row: ‘DAG-unreg’. $T = 1200$.

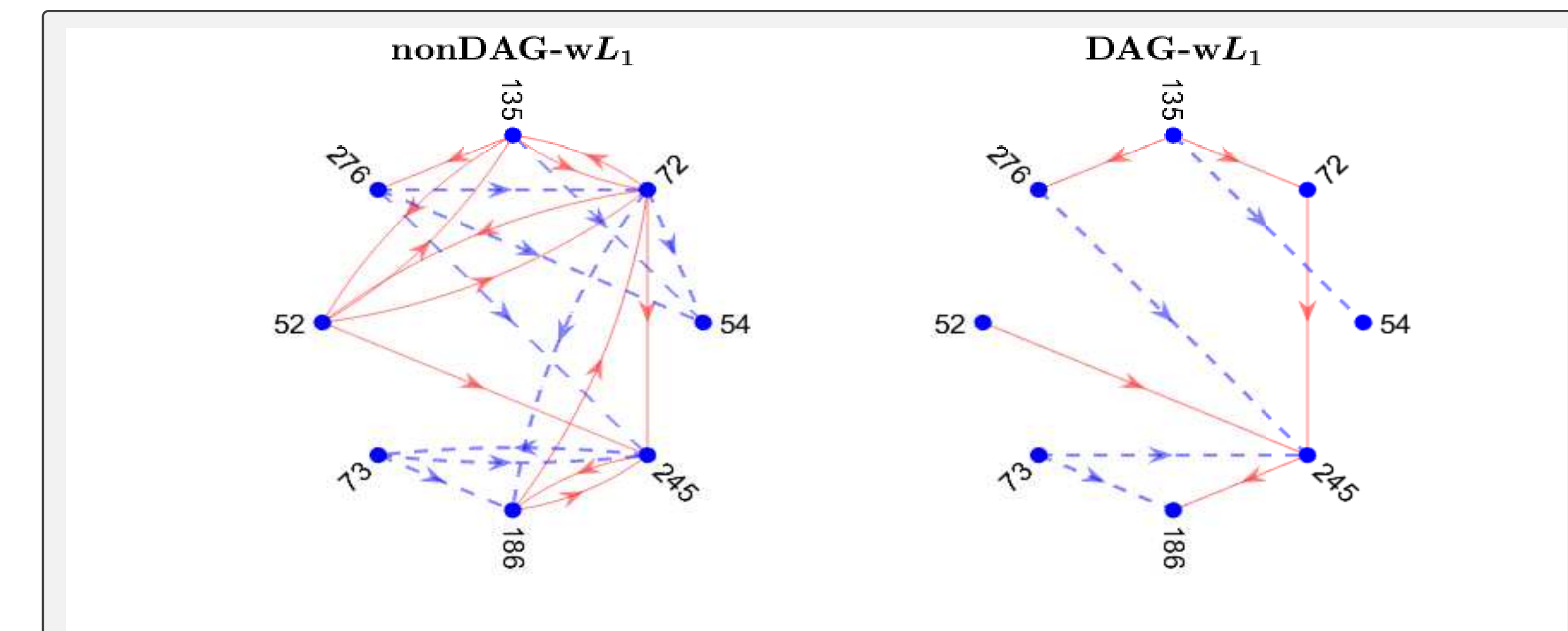


Figure: (Real spike train data: estimated sub-networks) Comparison of estimated sub-networks using two methods. Red solid lines with arrows: excitatory effects; blue dashed lines with arrows: inhibitory effects.

6. Conclusion and Discussion

- Flex-AL: faster DAG-constrained optimization with a global convergence guarantee.
- New insights into DAG constraints and sparsity regularization.

References:

- Gao, Zhang, & Zhou (2024). *IEEE Trans. Inf. Theory*.
- Zhang, Gao, & Jia (2024). *JMLR*.

