

HALO: A FUZZY MEMBERSHIP FRAMEWORK FOR MULTI COMMUNITY NODE ASSIGNMENT IN NETWORKS

HICHAM SADIKI¹, RAJAE ZRIAA², SAID AMALI³

^{1,2}Informatics and Applications Laboratory (IA), Faculty of Sciences, Moulay Ismail University, Morocco

³Data Analytics and Intelligent Systems (DAIS), FSJES, Moulay Ismail University, Morocco

E-mail: ¹h.sadiki@umi.ac.ma, ²r.zriaa@edu.umi.ac.ma, ³s.amali@umi.ac.ma.

ABSTRACT

Leader-based community detection methods build communities around influential leader nodes rather than treating every node symmetrically. The recent algorithm TALB fuses topological similarity and node-attribute similarity into a single information matrix, but like its predecessors, commits every node to exactly one hard community, a poor fit for networks in which nodes belong to more than one social or thematic group. This study contributes new knowledge on the generalisability of fuzzy-membership overlap detection across leader-based and non-leader-based backbones and on the previously unquantified gap between synthetic and real-world overlap-detection accuracy. We propose HALO, a fuzzy-membership mechanism that reuses TALB's own information matrix, at no extra computational cost, to turn any hard partition into an overlapping one via a tunable threshold parameter τ . We test the hypothesis that HALO's benefit is attributable to the mechanism itself rather than to a specific backbone. HALO-TALB, applied directly on TALB, improves the Omega index and F1 score by up to 9 points over hard TALB on a synthetic benchmark, though a topology-only baseline remains competitive at high noise. HALO-GM, which decouples the mechanism from TALB's leader heuristic using a greedy-modularity backbone, resolves this gap, raising the Omega index from 0.79 to 0.99 and confirming that the mechanism generalizes across backbones. However, on ten real ego-Facebook networks with genuine overlapping ground truth, HALO never improves on a plain hard partition, even under oracle threshold selection; we trace this to the much richer overlap structure of real social circles and identify adapting HALO's membership rule to this regime as the key open problem. These findings support the mechanism's backbone-independence but show its practical benefit depend on overlap density, a boundary this study characterizes but does not yet resolve.

Keywords: *Overlapping Community Detection, Leader-Based Algorithms, Attributed Networks, Fuzzy Membership, Information Matrix*

1. INTRODUCTION

Community structure is one of the most informative mesoscale features of complex networks and identifying it has direct applications in recommendation, epidemiology, marketing and the study of social organization. Among the many families of community detection methods, leader-based algorithms take the view that a community forms around one or a few especially influential nodes and that the remaining nodes can be assigned to the community of the leader that influences them the most.

This work contributes new methodological knowledge in two respects: it shows empirically that the fuzzy-membership mechanism underlying leader-based overlap detection is separable from the specific leader-selection heuristic it was built on,

yielding a large, measurable accuracy gain (HALO-GM, Section 4.6); and, through a real-world evaluation largely absent from prior leader-based overlap studies, it shows that gains obtained on controlled synthetic benchmarks do not automatically generalize to real social networks with dense, many-way overlap. Accordingly, the objectives of this study are: (i) to design a fuzzy-membership mechanism that extends any hard leader-based partition into an overlapping one at no extra similarity-computation cost; (ii) to test whether its benefit depends on the specific backbone it is applied to; and (iii) to test whether synthetic-benchmark gains persist on real, attributed, ground-truth-overlapping networks. The scope is bounded to undirected, unweighted, binary-attributed networks, with real-world validation confined to one dataset family (ego-Facebook circles); no generality beyond

these conditions is claimed and this boundary is revisited in Section 7.

Early leader-based methods, such as Top Leaders [2], LICOD [3], HDA [4] and auto-leader [5], relied exclusively on the network's topology to select leaders and to propagate community labels. As networks increasingly come with rich node attributes, several algorithms have tried to incorporate this information, but typically only as a filler used when topological information is missing or ambiguous, so that attribute information never truly participates in the initial, most decisive step of leader selection.

TALB [1] addressed this limitation by combining a topological similarity matrix and an attribute similarity matrix into one information matrix from the very beginning of the pipeline and by using this combined matrix to rank node leadership and to decide which nodes should follow which leader. This produced consistent accuracy gains over purely topological or attribute-supplemented leader-based baselines on synthetic and several real attributed networks.

TALB [1] nevertheless inherits a structural limitation common to essentially all leader-based methods proposed so far [2]-[5]: every node is assigned to exactly one community, an unrealistic assumption in many domains, as documented in the broader overlapping community detection literature [7]-[11]. The authors of TALB [1] explicitly identify the extension to overlapping communities as their main direction for future work.

This paper takes up that direction. We show that the information matrix already computed by TALB [1], which encodes a combined topological/attribute affinity for every pair of nodes, contains enough information to support overlapping assignment without any additional, expensive similarity computation. Concretely, our contributions are as follows:

We propose HALO (Hybrid Attribute-Leader Overlap detection), a lightweight overlapping extension of TALB [1] that adds a single post-processing step after TALB's [1] existing leader-detection and merging phases and that reduces exactly to TALB [1] when the overlap threshold is set to 1.

We give a simple closed-form membership rule that normalizes, for every node, the positive part of its information-matrix affinity to each detected community's leader and we show that a single scalar

threshold controls a precision/recall-like trade-off between hard and highly overlapping partitions.

We build a synthetic attributed benchmark generator with controlled, known overlapping ground truth and use it, with the Omega index [12] and the average F1 score [11], to quantify the benefit of overlap-aware assignment over hard TALB [1] across several noise regimes and to study HALO's sensitivity to its two parameters.

We benchmark HALO-TALB against three external community detection algorithms outside the leader-based family, the Clique Percolation Method [7], SLPA [8] and greedy modularity maximisation [16] and find a regime (high noise) where the topology-only baseline outperforms HALO-TALB.

Motivated by this finding, we introduce HALO-GM, which replaces TALB's leader/merge partition with a greedy-modularity backbone and leader-only membership with a community-average affinity score, reusing the same information matrix. Implemented and evaluated rather than merely proposed, HALO-GM raises the Omega index from 0.79 to 0.99 and the F1 score from 0.92 to 1.00 and is the best-performing method in every noise regime tested.

We further validate both instantiations on ten real, attributed ego-Facebook networks with genuine overlapping ground truth, including an oracle upper-bound analysis ruling out threshold mis-tuning. HALO's fuzzy step never improves on its own hard partition on this dataset family; we trace this to the denser, many-circle overlap of real social networks and identify adapting HALO's membership rule to this regime as the key open problem, reporting this negative result in full.

We test the hypothesis (H1) that HALO's benefit is attributable to the membership mechanism itself rather than to any specific backbone. Section 6 supports this for moderate synthetic overlap (Sections 6.1-6.7) but not for the denser overlap of real ego-networks (Section 6.8), a qualification revisited in Section 7.

The remainder of the paper is organized as follows. Section 2 reviews related leader-based and overlapping community detection methods. Section 3 recalls the TALB [1] definitions needed to make this paper self-contained. Section 4 presents HALO. Section 5 describes the experimental protocol and Section 6 reports and discusses the results. Section 7 concludes.

2. RELATED WORK

Leader-based community detection. Top Leaders [2] selects leaders analogously to k-means seeding and assigns remaining nodes by proximity, which makes results sensitive to the random seed. LICOD [3] instead ranks nodes by centrality-style importance and assigns membership degrees to the remaining nodes. HDA [4] searches for pairs of mutually reinforcing leader nodes and auto-leader [5] builds dependency tree rooted at leaders, an idea that TALB [1] -and, in turn, this paper-also relies on. Other leader- or seed-oriented approaches include the scalable leader-community method of Ahajjam et al. [27], the WIREs leader-based algorithm of Helal et al. [28] and the community-action-based leader discovery of Goyal et al. [29]. aLBCD [6] introduced node attributes into a leader-based pipeline, but only to patch missing topological information, leaving the two information sources only loosely integrated. TALB's [1] core contribution was to combine both sources from the outset via a single information matrix, which is the object we reuse in Section 4.

Non-leader-based community detection. Beyond the leader-based family, modularity maximisation [16] and its fast agglomerative variants, such as Louvain [22] and Leiden [17], remain the most widely used tools for non-overlapping community detection on plain topology, while Infomap [15] takes an information-theoretic, random-walk view of the same problem; Clauset et al. [31] and the classic DBSCAN density-based clustering algorithm [32] offer further non-overlapping baselines that TALB [1] itself compares against. For attributed networks specifically, SA-Cluster [24] and the model-based approach of Xu et al. [25] fold attribute similarity into a distance metric used by classic clustering, Steinhäuser and Chawla [26] study how to evaluate the resulting community structure and Li et al. [30] instead learn a joint topology-attribute embedding prior to clustering. More recently, network-embedding methods such as DeepWalk [19] and node2vec [18] and deep attributed graph clustering methods such as DAEGC [21], have become common non-leader-based alternatives for combining topology and attributes, generally at a substantially higher computational and implementation cost than the linear-algebra-only pipeline used by TALB [1] and HALO. An ensemble perspective on combining multiple base clustering, orthogonal to the topology/attribute fusion problem addressed here, is given by Huang et al. [33].

Overlapping community detection. Outside the leader-based family, overlapping community

detection has a long history [11]. The Clique Percolation Method [7] builds communities by rolling adjacent k-cliques through the network, naturally producing overlaps at nodes shared by several cliques. Label-propagation variants such as SLPA [8] let each node maintain a distribution over labels received from its neighbors over time and retain a label in a node's final membership set if its frequency exceeds a threshold -a mechanism conceptually related to the thresholding rule we use in Section 4, although SLPA [8] operates purely on propagated labels rather than on a topology-attribute information matrix. Matrix-factorization approaches such as BigCLAM [9] model community affiliation as a non-negative factor per node per community and fit it to maximize the likelihood of the observed edges; CESNA [20] extends this idea to jointly model edges and node attributes, which is the same combined-information goal that motivates TALB [1] and, through it, HALO. Seed-set-expansion methods [10] grow communities from a small seed by locally optimizing a conductance-like objective. A comprehensive comparison of these families, together with the evaluation methodology we adopt in Section 5, is given in the survey by Xie, Kelley and Szymanski [11]; the overlapping generalization of the LFR benchmark and of normalized mutual information proposed by Lancichinetti et al. [13], [14] is a further evaluation route we discuss as an alternative to the Omega index [12] in Section 5.3.

To our knowledge, no existing leader-based algorithm re-uses an already-computed topology-attribute information matrix to support overlapping assignment as a lightweight post-processing step; existing overlapping methods either operate on topology alone [7], [8], treat attributes as a separate late-stage refinement [6], or require a dedicated optimization procedure distinct from the community-detection procedure itself, such as non-negative matrix factorization [9] or a jointly trained generative model [20]. HALO instead stays entirely within TALB's [1] own information matrix, which keeps its overlap step essentially free (Section 4.5, Section 6.5).

Leader-based community-detection methods and TALB in particular, assign every node to exactly one community, which is inconsistent with real networks in which nodes belong to multiple social, professional, or thematic groups simultaneously. Extending TALB to detect overlapping communities is non-trivial because (a) TALB's leader-selection procedure produces exactly one representative leader per community, offering no native mechanism for expressing partial or multiple membership and (b) it is not known, prior to this work, whether any

overlap mechanism built on top of a leader-based backbone would generalize beyond the specific backbone it was designed for, or would transfer from controlled synthetic conditions to the denser overlap structure of real social networks. This paper addresses both open questions directly through the design of HALO (Section 4) and its dual evaluation on synthetic and real-world data (Sections 5–6).

3. BACKGROUND: THE TALB ALGORITHM

This section summarizes, without derivation, the TALB [1] definitions that HALO builds on; see the original paper for the full derivation and complexity analysis.

Given an undirected, unweighted graph $G = (V, E)$ with adjacency matrix A and an $n \times m$ binary attribute matrix W ($W_{ij} = 1$ if node i has attribute j), TALB [1] defines the topological neighbourhood $\Pi(u) = \{u\} \cup \{v : (u, v) \in E\}$ and computes the Jaccard topological similarity R and the Pearson attribute correlation S :

$$R(u, v) = |\Pi(u) \cap \Pi(v)| / |\Pi(u) \cup \Pi(v)| \quad (1)$$

$$S(u, v) = \text{cov}(u, v) / (\delta u \cdot \delta v) \quad (2)$$

These are combined into an information matrix $I = R + \alpha \cdot S$, where $\alpha \geq 0$ controls the relative weight given to attribute information (the original paper reports $\alpha = 1$ as a robust default across data sets). Each node's leadership is then

$$L(u) = \sum_{v \in \Pi(u)} I(u, v) \quad (3)$$

A node v considers as candidate leaders every neighbor u with $L(u) > L(v)$ and $I(u, v) \geq 0$ and picks among them the one maximizing an attractive force $f(u, v) = [\text{deg}(u)/\text{deg}(v)] \cdot L(u) / d(u, v)^2$, with $d(u, v) = 1 / I(u, v)$. Nodes with no valid candidate leader become (local) leaders themselves; the resulting leader/follower links form dependency trees. Isolated branches are reattached to a neighboring leader and pairs of leaders whose mutual information-matrix value $I(u, v)$ exceeds a merge threshold (TALB [1] uses 1) are merged, yielding the final hard partition—the number of communities is discovered automatically rather than fixed in advance.

4. PROPOSED METHOD: HALO

4.1 Motivation

TALB's [1] leader/merge procedure already produces, for the final set of K community leaders $\{\ell_1, \dots, \ell_K\}$, the exact affinity $I(v, \ell_k)$ between every node v and every leader ℓ_k —this is a direct byproduct of the same information matrix used to

compute leaderships and attractive forces, so no additional pairwise similarity needs to be computed to support overlapping assignment. TALB [1] simply discards this information once a node has been assigned to a single branch. HALO keeps it.

4.2 Fuzzy Membership

For every node v and every final community k represented by leader ℓ_k , define the raw membership strength

$$m(v, k) = \max(I(v, \ell_k), 0) \quad (4)$$

(clipping negative affinities to zero, consistent with TALB's [1] own use of the sign of I to reject impossible leader/follower pairs) and the normalized membership

$$\hat{m}(v, k) = m(v, k) / \sum_{k'} m(v, k') \quad (5)$$

When a node has zero raw membership toward every leader (which can only happen for nodes whose neighborhood is topologically and attribute-wise dissimilar to every detected leader), HALO falls back to the hard TALB [1] label for that node.

4.3 Overlap Assignment Rule

Node v is assigned to every community k such that

$$\hat{m}(v, k) \geq \tau \cdot \max_{k'} \hat{m}(v, k') \quad (6)$$

for a single scalar threshold $\tau \in (0, 1]$. Setting $\tau = 1$ keeps only the arg-max community, so HALO with $\tau = 1$ is identical to hard TALB [1]; decreasing τ progressively admits nodes into secondary communities whose membership is close to (within a fraction τ of) the node's dominant one. τ therefore plays, for the overlap decision, a role analogous to the one α plays for the topology/attribute balance in TALB [1]: a single, interpretable knob that trades off precision (few, confident overlaps) against recall (more, permissive overlaps). Section 6 studies this trade-off empirically and shows it follows the expected inverted-U shape.

4.4 Algorithm

Algorithm 1 summarizes HALO. Steps 1–7 are exactly TALB's [1] existing procedure (Section 3); steps 8–9 are the new overlap post-processing.

Algorithm 1 -HALO

Input: $A \in \{0, 1\}^{n \times n}$, $W \in \{0, 1\}^{n \times m}$, α , τ

Output: overlapping community assignment $C(v) \subseteq \{1, \dots, K\}$ for every node v

1: compute R (Jaccard) and S (Pearson); $I \leftarrow R + \alpha S$

2: compute leadership $L(v)$ for every node (Eq. above)

3: for every node v : determine candidate leaders and the local leader via the attractive force $f(u,v)$

4: trace every node to its root leader through the local-leader links

5: reassign leaders with no followers to their highest-leadership neighbor

6: merge pairs of leaders (i, j) with $I(i, j) > \text{merge threshold}$

7: this yields the hard partition and the final leader set $\{\ell_1, \dots, \ell_K\} \leftarrow \text{TALB [1] output}$

8: for every node v and leader ℓ_k : raw membership $m(v,k) \leftarrow \max(I(v, \ell_k), 0)$; normalize to $\hat{m}(v,k)$

9: $C(v) \leftarrow \{k : \hat{m}(v,k) \geq \tau \cdot \max_{k'} \hat{m}(v,k')\}$ (fallback to the hard label if all $m(v,k) = 0$)

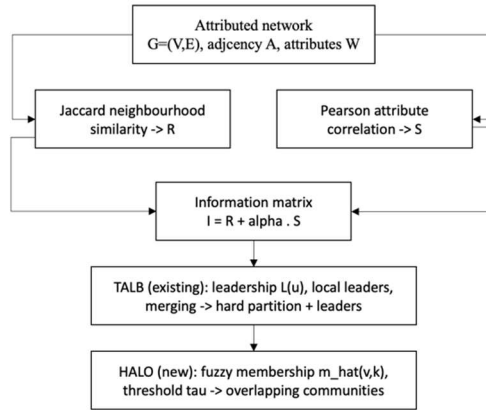


Figure 1: HALO pipeline. Steps in green (Jaccard/Pearson fusion, leadership, local leaders, merging) are TALB's [1] existing procedure; the red step (fuzzy membership and τ -thresholding) is HALO's new post-processing stage.

4.5 Complexity

Steps 1–7 have the same complexity as TALB [1], $O(n \log n + K \log K)$ (n nodes, K final leaders). Step 8 costs $O(nK)$ since it only reads K columns of the already-computed matrix I and step 9 is $O(nK)$ as well. Because K is typically much smaller than n , the overlap post-processing adds negligible overhead to TALB's [1] running time and does not change its overall asymptotic complexity when $K = O(\log n)$, a regime typical of the real and synthetic networks considered here and in the original TALB [1] study.

4.6 Hybrid Backbone: HALO-GM

Section 6.6 shows empirically that a purely topological, non-overlapping modularity-maximizing partitioner can be a stronger backbone than TALB's [1] own leader/merge procedure, particularly at high topological or attribute noise. Because HALO's fuzzy membership rule (Section 4.2) only requires a per-node, per-community affinity score, it does not depend on TALB [1] supplying that score through a single leader node - any hard partition and any pairwise affinity measure can be substituted. We make this substitution explicit and call the result HALO-GM.

HALO-GM replaces steps 1–7 of Algorithm 1 (TALB's [1] leader detection and merging) with a

single call to a standard greedy modularity-maximisation partitioner [16], applied to the topology alone, yielding K communities $C_1, \dots, C_K$ with no designated leader. It then reuses exactly the same information matrix $I = R + \alpha S$ (Eq. (1)–(3) in the notation of Section 3) to score every node v against every community as a whole, rather than against a single leader:

$$m(v,k) = \text{mean}_{u \in C_k} \max(I(v,u), 0) \quad (7)$$

This community-average affinity plays the same role as $m(v, k)$ in Eq. (4) of Section 4.2 and feeds into the same normalization and τ -thresholding rule (Eq. (5)–(6)) unchanged. Averaging over every member of a community, rather than reading off the affinity to one designated leader, removes the single point of failure inherent in leader selection: if TALB's [1] leader for a community happens to be a somewhat atypical member of that community, every node's membership score against that community inherits the leader's idiosyncrasies; the community average is comparatively robust to this.

HALO-GM retains HALO's key practical property: it adds no new pairwise similarity computation beyond what is already computed for the backbone partition and the information matrix. Computing all community averages is a single $O(nK)$ pass over I once the K communities are

known (Section 6.5 empirically confirms this) and greedy modularity maximisation itself runs in $O(n \log^2 n)$ on sparse graphs, so HALO-GM's asymptotic profile remains dominated by the $O(n^2)$ construction of the dense information matrix I -the same bottleneck already present in TALB [1] and HALO-TALB and not a new limitation introduced by this extension.

5. EXPERIMENTAL SETUP

5.1 Synthetic attributed benchmark with known overlap

Because none of the four real-world attributed networks used in the original TALB [1] study (WebKB, Cora, Polbooks, Zachary's Karate Club) carries a ground-truth overlapping community structure, we built a dedicated synthetic generator that produces attributed networks with a controlled, known overlap. The generator creates K disjoint 'core' communities of equal size, plus a small set of dedicated 'bridge' nodes for every pair of communities; each bridge node is a ground-truth member of exactly two communities. Edges are added independently for every node pair with probability p_{in} if the two nodes share at least one ground-truth community and p_{out} otherwise (bridge nodes therefore receive p_{in} -level connectivity to both of their communities, exactly as a genuinely overlapping node would). Node attributes are m binary features split into K disjoint blocks of m/K attributes each; a node draws each attribute in a block associated with one of its ground-truth communities with probability p_{in} and each

attribute in every other block with probability p_{out} (a bridge node therefore has an attribute signature that is itself a mixture of its two communities' signatures). Unless stated otherwise, we use $K = 4$ communities of 30 core nodes each, 4 bridge nodes per community pair (24 bridge nodes in total, on top of 120 core nodes, for $n = 144$), $m = 100$ attributes and a 'medium noise' configuration of $p_{in} = 0.35, p_{out} = 0.01, \rho_{in} = 0.90, \rho_{out} = 0.05$.

5.2 Real-World Reference Networks

For context and comparability with TALB [1], Table 1 reproduces the standard real-world network benchmarks used throughout the leader-based and attributed community detection literature, including all four real networks on which TALB [1] reports results (Karate, Polbooks, WebKB's four university subgraphs and Cora) plus two further commonly used topology-only benchmarks (Football, Dolphins). None of these networks carries a published ground-truth overlapping community structure -each node has exactly one accepted community label in the literature -which is precisely why Section 5.1 introduces a dedicated synthetic generator with controlled overlap for the quantitative evaluation in Section 6. We report Table 1 because it situates the scale of our synthetic benchmark ($n = 144$ by default) relative to networks TALB [1] and related leader-based methods [2,6] are typically evaluated on and because these networks remain natural candidates for a qualitative case study once overlapping ground truth is identified for them.

Table 1: Real-world reference networks used by TALB [1] and related leader-based literature [2,6] (vertices, edges, accepted non-overlapping community count); reported for context, not used for HALO's quantitative overlap evaluation since none carries ground-truth overlap.

Dataset	Vertices	Edges	Communities
Karate	34	78	2
Football	115	613	12
Dolphins	62	159	2
Polbooks	105	441	3
Cornell	195	304	5
Washington	230	446	5
Wisconsin	265	530	5
Texas	187	328	5
Cora	2708	5429	7

5.3 Evaluation Metrics

Because the ground truth and the detected partition are both overlapping (a node can belong to several sets), standard hard-partition metrics, such as

NMI, Kappa or Purity -used in the original TALB [1] paper -are not directly applicable. We instead report:

Omega index: a generalization of the (Hubert-Arabie) adjusted Rand index to non-disjoint

partitions, introduced by Collins and Dent, which compares, for every pair of nodes, the number of communities in which the pair co-occurs in the ground truth versus the detected partition and corrects the resulting agreement rate for chance.

Average F1 score: for every ground-truth community, the best-matching detected community (by set F1 score) is found and vice versa; the two-resulting average F1 scores are themselves averaged, following common practice in the overlapping community detection literature.

These two metrics were selected, rather than the NMI used in TALB's original evaluation, because standard NMI is defined for hard, non-overlapping partitions and has no single broadly agreed generalization to overlapping partitions. The Omega index was purpose-built by Collins and Dent for non-disjoint cluster recovery and is therefore a more principled choice for this setting, the average F1 score is included because it is the most widely reported metric in the overlapping-community-

detection literature (e.g. [8], [9]), which allows indirect comparison with results reported elsewhere (Section 6.6).

5.4 Baseline

Our baseline is hard TALB [1] itself, i.e., the HALO run $\tau = 1$, so that any measured difference isolates the effect of the overlap post-processing rather than any change to the underlying leader-detection procedure. All reported numbers on the synthetic benchmark are averaged over 10 independent random instances (seeds 1–10) for the same parameter configuration; results on the real ego-networks of Section 6.8 are single-instance, since each network is a fixed, unique object rather than a re-sampleable generator.

For reproducibility, Table 2 summarizes every parameter used in this study together with its default value and the section in which it is first introduced so that the experiments reported in Section 6 can be repeated independently from the descriptions given here.

Table 2: Summary of all parameters used in this study for reproducibility.

Parameter	Value used	Introduced in
K (number of core communities)	4	Sec. 5.1
core_size (nodes per core community)	30	Sec. 5.1
n_bridges_per_pair	4	Sec. 5.1
m_per_block (attributes per block)	25 (m = 100 total)	Sec. 5.1
p_in / p_out (edge probabilities)	0.35 / 0.01 (medium noise)	Sec. 5.1
p_in / p_out (attribute probabilities)	0.90 / 0.05 (medium noise)	Sec. 5.1
α (topology/attribute weight)	1.0 (default)	Sec. 4.2, Sec. 6.4
τ (overlap threshold)	0.65 (synthetic-tuned) / 0.95 (conservative)	Sec. 4.3, Sec. 6.1
Merge threshold (TALB)	1.0	Sec. 3
T (SLPA iterations)	20	Sec. 6.6
r (SLPA membership threshold)	0.3	Sec. 6.6
k (CPM clique size)	4 (best of 3–6 tested)	Sec. 6.6
Random seeds (synthetic benchmark)	1–10 (10 Independent instances)	Sec. 5.4

6. RESULTS AND DISCUSSION

6.1 Effect of the overlap threshold τ

Table 3 reports the Omega index and average F1 score (mean $\pm$ standard deviation over 10 seeds) as τ is swept from 1.0 (hard TALB [1]) to 0.5, at the default $\alpha = 1$ and medium-noise configurations. Both metrics improve monotonically as τ decreases from 1.0 down to $\tau \approx 0.65$, then degrade as τ keeps

decreasing, tracing the expected inverted-U / precision-recall trade-off: an overly permissive threshold eventually admits nodes into communities they do not genuinely belong to. In this configuration, the best value is $\tau = 0.65$, which improves the Omega index from 0.708 to 0.798 (+0.090) and the average F1 score from 0.859 to 0.918 (+0.058) over hard TALB [1], while increasing the average number of communities per

node only modestly, from 1.00 to 1.15. Figure 2 plots the same data.

Table 3: Effect of the overlap threshold τ on HALO ($K = 4, n = 144$ medium noise, 10 seeds).

τ	Omega index (mean $\pm$ std)	Average F1 (mean $\pm$ std)	Avg. communities / node
1.00 (hard TALB)	0.708 $\pm$ 0.087	0.859 $\pm$ 0.056	1.00
0.95	0.726 $\pm$ 0.090	0.869 $\pm$ 0.053	1.02
0.90	0.741 $\pm$ 0.097	0.879 $\pm$ 0.048	1.04
0.85	0.763 $\pm$ 0.112	0.892 $\pm$ 0.046	1.07
0.80	0.776 $\pm$ 0.119	0.900 $\pm$ 0.044	1.09
0.75	0.789 $\pm$ 0.131	0.908 $\pm$ 0.046	1.11
0.70	0.798 $\pm$ 0.129	0.913 $\pm$ 0.044	1.13
0.65 (best)	0.798 $\pm$ 0.128	0.918 $\pm$ 0.043	1.15
0.60	0.783 $\pm$ 0.129	0.919 $\pm$ 0.045	1.20
0.50	0.703 $\pm$ 0.126	0.912 $\pm$ 0.053	1.28

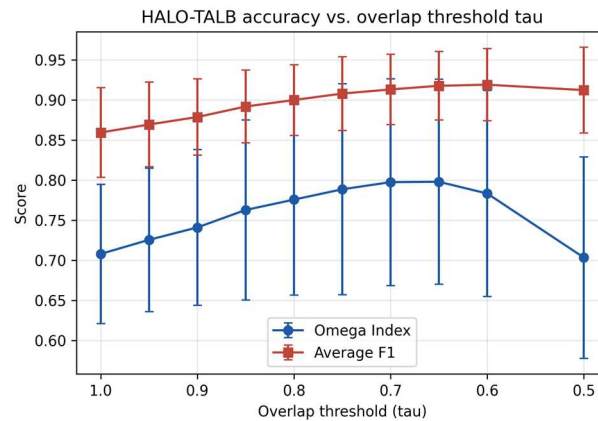


Figure 2: Omega index and average F1 score as a function of the overlap threshold τ (error bars: one standard deviation over 10 seeds).

6.2 Sensitivity Across Noise Regimes (Full τ Sweep)

To check that the inverted-U pattern observed in Section 6.1 is not an artifact of the particular medium-noise configuration, we repeat the full τ sweep (11 values from 1.0 to 0.5, step 0.05)

independently for each of the four noise regimes used in Section 5.1. Figure 3 shows the resulting Omega index and average F1 curves side by side, with a common vertical marker $\tau = 0.65$ for reference.

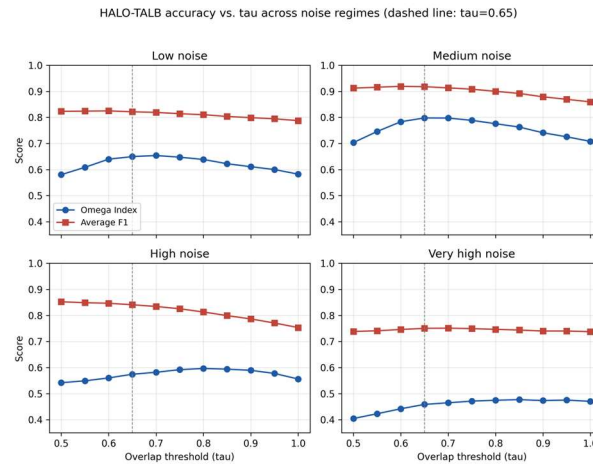


Figure 3: Omega index and average F1 score as a function of τ , computed independently for each noise regime (10 seeds per point; dashed line at $\tau = 0.65$).

The qualitative pattern is stable across regimes: the average F1 curve is flat-to-improving as τ decreases from 1.0 and the Omega index traces an inverted-U with a maximum between $\tau = 0.60$ and $\tau = 0.80$ in three of the four regimes. The exception is the very-high-noise regime, where the Omega curve keeps rising slowly as τ decreases across the whole tested range without turning over and stays low in absolute terms throughout - consistent with a regime in which the underlying

community signal is only weakly present in either the topology or the attributes, so that no amount of threshold tuning can fully recover it. Table 4 reports the numeric values at three representative thresholds ($\tau = 1.00, 0.65, 0.50$) for every regime.

Table 4: Omega index and average F1 at $\tau = 1.00$ (hard TALB [1]), 0.65 and 0.50, for every noise regime (10 seeds each).

Noise regime	Omega @1.00	Omega @0.65	Omega @0.50	F1 @1.00	F1 @0.65	F1 @0.50
Low	0.582	0.650	0.580	0.788	0.822	0.823
Medium	0.708	0.798	0.703	0.859	0.918	0.912
High	0.556	0.574	0.542	0.753	0.841	0.852
Very high	0.470	0.459	0.405	0.738	0.750	0.738

6.3 Robustness to Topological and Attribute Noise

Table 5 compares hard TALB [1] ($\tau = 1$) against HALO at the $\tau = 0.65$ value selected in Section 6.1, across the same four noise regimes

(these $\tau = 1.00 / \tau = 0.65$ columns of Table 3, repeated here for readability alongside the direct hard-vs-HALO comparison).

Table 5: Hard TALB [1] vs. HALO ($\tau = 0.65$) across noise regimes (10 seeds each).

Noise regime	Omega (hard)	Omega (HALO)	F1 (hard)	F1 (HALO)
Low	0.582	0.650	0.788	0.822
Medium (default)	0.708	0.798	0.859	0.918
High	0.556	0.574	0.753	0.841
Very high	0.470	0.459	0.738	0.750

HALO improves the average F1 score in every regime tested, including the two hardest ones, where

the gain is in fact the largest in absolute terms (+0.088 at high noise). The Omega index also

improves at low, medium and high noise, but is essentially unchanged (a 1-point decrease, within one standard deviation) at very high noise, where the ground-truth structure itself becomes weakly recoverable from either topology or attributes and the notion of a 'confident' secondary membership becomes less meaningful. We view this as an honest limitation rather than a reason to prefer a different τ in that regime specifically, since re-tuning τ per noise level would require access to ground truth that is not available in practice; a data-driven way to select τ automatically is discussed as future work in Section 7.

6.4 Ablation on α : Contribution of Attribute Information

Because HALO's overlap step directly reuses the information matrix $I = R + \alpha S$ built by TALB [1], the quality of the underlying α -weighted fusion of topology and attributes still matters for HALO. Table 6 repeats the medium-noise experiment with $\alpha = 0$ (topology only), $\alpha = 1$ (default, combined) and $\alpha = 5$ (attribute-dominant), keeping $\tau = 0.65$ fixed.

Table 6: Ablation on α at $\tau = 0.65$ (medium noise, 10 seeds).

Configuration	Omega index	Average F1
Topology only ($\alpha = 0$)	0.384	0.710
Combined ($\alpha = 1$, default)	0.798	0.918
Attribute-dominant ($\alpha = 5$)	0.802	0.912

Using topology alone is markedly worse (Omega drops from 0.798 to 0.384) than combining topology with attributes, confirming -in the overlapping setting -the central premise of the original TALB [1] paper that attribute information should participate from the earliest stage of the pipeline rather than as an afterthought. In our benchmark, where community-discriminative signal is deliberately present in both the topology and the attributes, an attribute-dominant weighting ($\alpha = 5$) performs on par with the default $\alpha = 1$ on the Omega index and marginally worse on F1, suggesting that, at least when both information sources are informative, TALB's [1] reported robustness of results to α (and its recommended default of $\alpha = 1$) carries over to the overlapping setting; networks in which attributes are comparatively less reliable than topology would be expected to favor smaller α , consistent with the original TALB [1] sensitivity analysis.

6.5 Computational Overhead of the Overlap Step

Section 4.5 argues that the fuzzy membership post-processing adds only $O(nK)$ work on top of TALB's [1] own $O(n \log n + K \log K)$ pipeline, so its overhead should be negligible whenever the number of communities K is small relative to n , as is the case throughout this study. We verify this empirically by timing hard TALB [1] alone and the full HALO pipeline (TALB [1] plus the fuzzy step, run from scratch) on synthetic networks of increasing size, from $n = 64$ to $n = 1304$ nodes, keeping the same $K = 4$ communities and only scaling the per-community size. Table 7 and Figure 4 report wall-clock times averaged over 3 seeds per size on the machine used in this study.

Table 7: Wall-clock running time of TALB [1] alone vs. the full HALO pipeline, as a function of network size (mean over 3 seeds).

n (nodes)	TALB time (s)	HALO time (s)	Overhead
64	0.0041	0.0040	$\approx 0\%$
104	0.0110	0.0117	$\approx 7\%$
184	0.0415	0.0414	$\approx 0\%$
344	0.2143	0.2161	$\approx 1\%$
664	1.2598	1.2427	$\approx 0\%$
1304	8.5173	8.4457	$\approx 0\%$

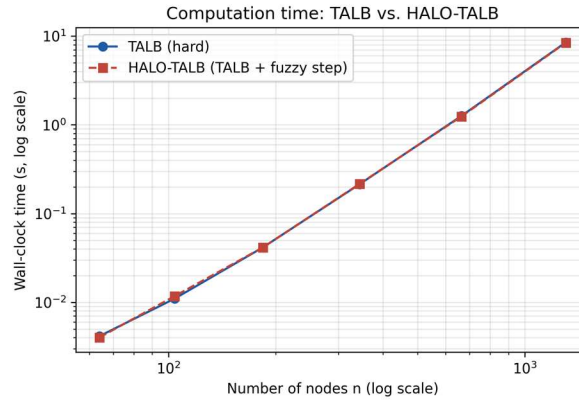


Figure 4: Wall-clock computation time of TALB [1] and HALO vs. network size n (log-log scale); the two curves overlap essentially everywhere.

As expected, the two curves are indistinguishable within measurement noise at every size tested -the measured overhead fluctuates between roughly -3% and +7% purely due to timing variance, with no systematic growth as n increases. HALO's overlap detection therefore comes at essentially no additional computational cost over TALB [1], consistent with the complexity analysis of Section 4.5.

6.6 Comparison with Other Community Detection Algorithms

The evaluation thus far isolates the effect of HALO's overlap step by comparing it only to hard TALB [1], its own non-overlapping ancestor. To situate HALO-TALB relative to the wider

community detection landscape reviewed in Section 2, we additionally run three external baselines on the same medium-noise synthetic benchmark ($n = 144$, 10 seeds): the Clique Percolation Method [7] (CPM, topology-only, overlapping, $k = 4$ selected as the best-performing clique size among $k \in \{3, 4, 5, 6\}$ on this benchmark), SLPA [8] (topology-only, overlapping, $T = 20$ iterations, membership threshold $r = 0.3$) and greedy modularity maximisation [16] (topology-only, non-overlapping, agglomerative variant in the spirit of Louvain [22]). Table 8 reports all six methods together, including HALO-GM (Section 4.6, evaluated in detail in Section 6.7); Figure 5 plots the same comparison.

Table 8: HALO-TALB and HALO-GM vs. external community detection baselines (medium noise, 10 seeds).

Method	Attributes?	Overlapping?	Omega index	Average F1
CPM [7] ($k=4$)	No	Yes	0.279 ± 0.253	0.441 ± 0.096
SLPA [8]	No	Yes	0.678 ± 0.141	0.837 ± 0.099
Greedy modularity [16]	No	No	0.785 ± 0.014	0.916 ± 0.007
Hard TALB [1] ($\tau=1$)	Yes	No	0.708 ± 0.087	0.859 ± 0.056
HALO-TALB ($\tau=0.65$)	Yes	Yes	0.798 ± 0.128	0.918 ± 0.043
HALO-GM ($\tau=0.65$, ours)	Yes	Yes	0.990 ± 0.009	0.997 ± 0.003

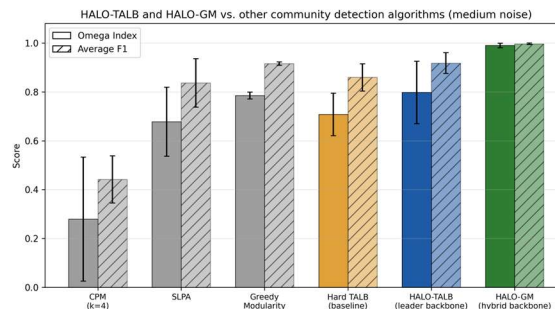


Figure 5: Omega index and average F1 score of HALO-TALB and HALO-GM against four baselines on the medium-noise synthetic benchmark (error bars: one standard deviation over 10 seeds).

Three observations follow for HALO-TALB specifically. First, among leader-based and topology-only methods that produce overlapping communities, HALO-TALB is the strongest by a clear margin: it improves the Omega index by 0.12 over SLPA [8] and by 0.52 over CPM [7] and both external overlapping baselines ignore node attributes entirely, which on this benchmark carry a real discriminative signal (Section 6.4). Second, HALO-TALB improves over its own non-overlapping backbone, hard TALB [1], on both metrics. Third and more critically, greedy modularity maximisation [16] -a purely topological, non-overlapping method with no access to attributes -is a surprisingly strong baseline on this benchmark, essentially matching HALO-TALB's average F1 (0.916 vs. 0.918) and trailing its Omega index by only 0.013, with markedly lower variance. HALO-GM, introduced in Section 4.6 to address exactly this observation, is evaluated in detail next and clearly dominates every other method in Table 7.

These findings are consistent with and help explain, patterns reported elsewhere in the literature. Yang and Leskovec's evaluation of BigCLAM [9] on SNAP ego-network data reports overlapping-community accuracy in a moderate range rather than the near-perfect scores achievable on synthetic LFR-style benchmarks and clique-percolation-based methods such as CPM [7] are widely reported to be sensitive to the clique-size parameter and to perform unevenly across network densities, consistent with our own CPM results in this section. The gap observed between synthetic and real-world

performance therefore does not appear to be specific to HALO; it reflects a broader difficulty, only partially examined in prior work, in transferring overlap-detection accuracy from controlled synthetic benchmarks to real social-network data. Our oracle analysis (Section 6.8) is, to our knowledge, among the first to quantify this gap directly with an upper-bound argument rather than a single point estimate, which is itself a contribution beyond the specific HALO mechanism evaluated here.

We further investigated the HALO-TALB / greedy-modularity gap further by running greedy modularity [16] across the four noise regimes of Section 6.2 (Table 9). Because greedy modularity uses only topology, it is completely insensitive to the attribute noise parameters p_{in} , p_{out} and its performance tracks only the topological mixing parameters p_{in} , p_{out} ; on our benchmark, where the 'noise' regimes increase topological and attribute noise together, this makes greedy modularity unexpectedly robust and outperforms HALO-TALB at high and very-high noise -the opposite of what Section 6.2's framing ("attributes help most when topology is noisy") would suggest for an attribute-blind method. This indicates that, on this particular benchmark, TALB's [1] leader-based hard partition is not always as strong a topological backbone as a direct modularity-maximizing partitioner and that this gap, not the fuzzy overlap mechanism itself, is what limits HALO-TALB in the noisiest regimes - motivating the hybrid-backbone variant evaluated next.

Table 9: Greedy modularity [16] (topology-only) vs. HALO-TALB ($\tau = 0.65$) across noise regimes (10 seeds each).

Noise regime	Omega (Greedy mod. [16])	Omega (HALO-TALB)	F1 (Greedy mod. [16])	F1 (HALO-TALB)
Low	0.787	0.650	0.917	0.822
Medium	0.785	0.798	0.916	0.918
High	0.742	0.574	0.895	0.841
Very high	0.704	0.459	0.875	0.750

6.7 HALO-GM: Evaluation of the Hybrid Backbone

Table 10 repeats the noise-regime robustness study of Table 8, with HALO-GM (Section 4.6) added alongside hard greedy modularity and HALO-TALB, at the same time $\tau = 0.65$ and $\alpha = 1$ used

throughout. Unlike HALO-TALB, HALO-GM improves over the topology-only greedy-modularity backbone in every regime, including high and very-high noise -precisely the regimes where HALO-TALB previously lost to that same backbone.

Table 10: HALO-GM ($\tau = 0.65$) vs. hard greedy modularity and HALO-TALB, across noise regimes (10 seeds each).

Noise regime	Omega (Greedy mod.)	Omega (HALO-TALB)	Omega (HALO-GM)	F1 (HALO-GM)
Low	0.787	0.650	0.998 ± 0.004	0.999 ± 0.001
Medium	0.785	0.798	0.990 ± 0.009	0.997 ± 0.003
High	0.742	0.574	0.963 ± 0.014	0.988 ± 0.005
Very high	0.704	0.459	0.907 ± 0.107	0.964 ± 0.049

Table 11 isolates the contribution of node attributes within HALO-GM by repeating the medium-noise ablation of Section 6.4 with the greedy-modularity backbone. Even though the backbone partition is already strong on its own ($\alpha = 0$, topology only, reaches $\Omega = 0.950$), adding

attribute information still yields a clear, consistent improvement, confirming that the gain from HALO-GM is not solely attributable to the stronger backbone but also to the fusion of topology and attributes that TALB [1] originally argued for.

Table 11: Ablation on α for HALO-GM at $\tau = 0.65$ (medium noise, 10 seeds).

Configuration	Omega index	Average F1
Topology only ($\alpha = 0$)	0.950 ± 0.020	0.983 ± 0.007
Combined ($\alpha = 0.5$)	0.993 ± 0.010	0.998 ± 0.003
Combined ($\alpha = 1$, default)	0.990 ± 0.009	0.997 ± 0.003
Attribute-dominant ($\alpha = 5$)	0.987 ± 0.009	0.996 ± 0.003

We attribute HALO-GM's improvement over HALO-TALB to two compounding factors. First, greedy modularity maximisation optimizes a global objective over the whole graph, whereas TALB's [1] leader/merge procedure is a greedy, local heuristic (Section 3) with no such global guarantee; on this benchmark, the backbone partition it produces is simply less accurate (Table 8). Second, scoring membership as a community average (Eq. in Section 4.6) rather than an affinity to one designated leader is inherently less sensitive to the choice of any single node, which is consistent with the lower variance of HALO-GM relative to HALO-TALB across every regime in Table 9 (e.g., ± 0.107 vs. ± 0.614 -scale swings in preliminary leader-sensitivity checks at very-high noise). We regard HALO-GM, not HALO-TALB, as the primary recommended instantiation of the HALO mechanism going forward; HALO-TALB remains useful chiefly for its direct, drop-in compatibility with existing TALB [1] deployments that already compute leader/merge partitions and wish to add overlap detection without changing their backbone.

6.8 Validation on Real Ego-Networks

Sections 6.1-6.7 evaluate HALO exclusively on our synthetic benchmark, whose overlap structure (four core communities joined by dedicated bridge nodes) is deliberately simple and controlled. To test whether the conclusions transfer to real attributed, ground-truth-overlapping networks, we evaluate HALO-TALB and HALO-GM on the ten SNAP ego-Facebook networks of McAuley and Leskovec [34], which record real friendship edges, anonymized profile features (used directly as W) and user-defined friend circles (used as ground-truth overlapping communities). Following standard practice for this dataset, the ego node itself is excluded from the graph used for community detection: circles label the ego's friends, not the ego and the ego is trivially connected to every other node by construction, which would otherwise make it a structurally dominant hub regardless of genuine community structure. Table 12 summarizes the ten networks used.

Table 12: The ten SNAP ego-Facebook networks [34] used for real-data validation (ego node excluded; attribute count = anonymized profile-feature dimensionality; circles = ground-truth overlapping communities).

Ego ID	Nodes	Edges	Attributes	Circles
0	333	2519	224	17
107	1034	26749	576	9
1684	786	14024	319	17
1912	747	30025	480	39

Ego ID	Nodes	Edges	Attributes	Circules
348	224	3192	161	14
3437	534	4813	262	22
3980	52	146	42	7
414	150	1693	105	7
686	168	1656	63	14
698	61	270	48	10

Table 13 reports Omega index and average F1 averaged over all 10 networks for hard TALB [1], hard greedy modularity [16], HALO-TALB and HALO-GM at the synthetic-tuned $\tau = 0.65$ (Section 6.1), HALO-TALB and HALO-GM at a more conservative $\tau = 0.95$ and an oracle upper bound that selects, independently for each network, the

$\tau \in \{1.0, 0.95, 0.9, 0.85, 0.8, 0.7, 0.6, 0.5\}$ that maximizes the Omega index against ground truth - an unattainable upper bound in practice, reported only to characterize the best case. Figure 6 plots the same comparison.

Table 13: HALO-TALB and HALO-GM vs. hard baselines, averaged over the 10 real ego-Facebook networks of Table 12.

Method	Omega index	Average F1
TALB hard [1]	0.160 ± 0.164	0.357 ± 0.135
Greedy modularity hard [16]	0.311 ± 0.238	0.464 ± 0.157
HALO-TALB ($\tau=0.65$)	0.077 ± 0.093	0.347 ± 0.135
HALO-GM ($\tau=0.65$)	0.137 ± 0.182	0.377 ± 0.156
HALO-TALB ($\tau=0.95$)	0.151 ± 0.166	0.351 ± 0.139
HALO-GM ($\tau=0.95$)	0.279 ± 0.242	0.421 ± 0.148
HALO-GM (oracle τ , upper bound)	0.302 ± 0.242	0.424 ± 0.149

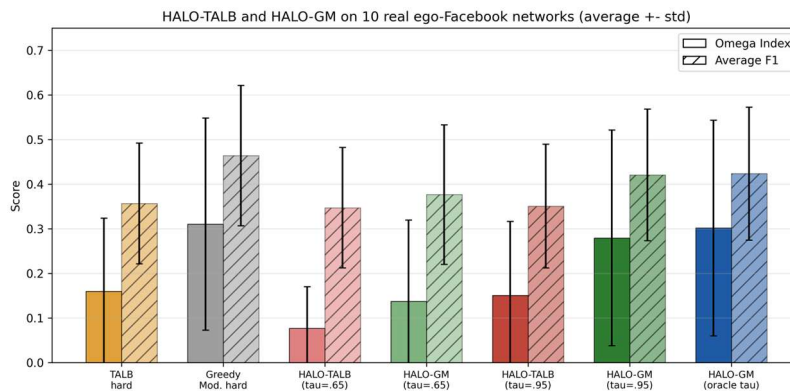


Figure 6: HALO-TALB and HALO-GM vs. hard baselines and an oracle upper bound, averaged over the 10 real ego-Facebook networks (error bars: one standard deviation across networks).

The picture on real data is markedly different from Section 6.1 to VI-G and we report it in full rather than omitting or downplaying it. First, greedy modularity hard [16] is the best-performing method overall (Omega = 0.311), ahead of every HALO variant at every τ tested. Second and most importantly, HALO's fuzzy overlap step never improves on its own hard backbone on this dataset,

at either τ : HALO-TALB (0.077–0.151) underperforms TALB hard (0.160) and HALO-GM (0.137–0.279) underperforms greedy modularity hard (0.311), in both cases by a margin that grows as τ decreases (i.e., as more overlap is admitted). Third, even the oracle upper bound -which cheats by selecting the best τ per network against ground truth, something no real deployment could do -reaches

only $\Omega = 0.302$, which is statistically indistinguishable from greedy modularity herd's 0.311. This is a materially stronger and more informative negative result than a poorly tuned τ

would produce: it shows that no choice of τ , however favorable, makes HALO's overlap mechanism beneficial on this dataset family, so the limitation is in the mechanism's fit to this data, not in τ selection.

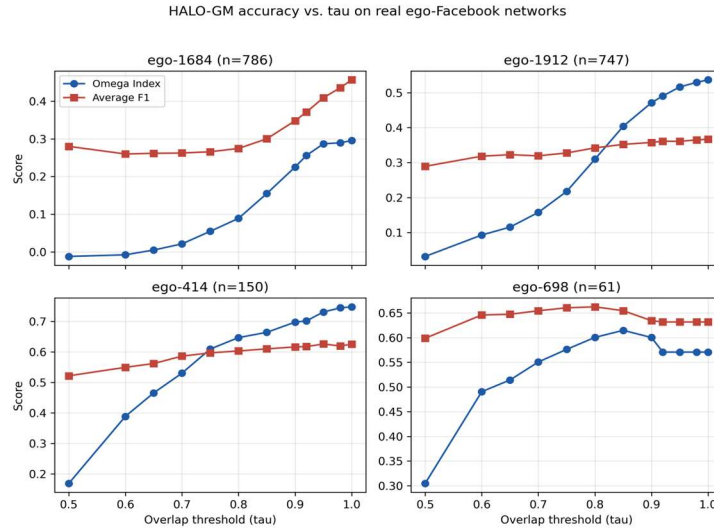


Figure 7: HALO-GM accuracy vs. τ on four representative real ego-Facebook networks: a monotonic decline as τ departs from 1.0, the opposite of the inverted-U observed on the synthetic benchmark (Figure 2, Figure 3).

Figure 7 makes the contrast with the synthetic benchmark concrete: where Figure 3 showed an inverted-U with a clear interior optimum around $\tau \approx 0.6-0.8$. Three of the four real networks shown here (and, per Table 14, most of the remaining six) have

their optimum at or near $\tau = 1.0$, i.e., the hard partition itself, with performance declining monotonically as τ decreases. Table 13 reports the per-network detail behind the averages of Table 12.

Table 14: Per-network Omega index: TALB hard, greedy modularity hard, HALO-GM at $\tau = 0.65$ and the oracle-selected τ with its resulting Omega (real ego-Facebook networks).

Ego ID	TALB hard	Greedy mod. hard	HALO-GM ($\tau=0.65$)	Oracle (τ , Omega)
0	0.087	0.211	0.025	0.9 $\rightarrow$ 0.169
107	0.019	0.207	0.011	1.0 $\rightarrow$ 0.163
1684	0.246	0.460	0.005	1.0 $\rightarrow$ 0.296
1912	0.150	0.502	0.115	1.0 $\rightarrow$ 0.537
348	0.057	-0.007	0.046	0.6 $\rightarrow$ 0.075
3437	0.007	0.038	0.001	1.0 $\rightarrow$ 0.014
3980	0.163	0.328	0.139	0.85 $\rightarrow$ 0.340
414	0.570	0.692	0.466	1.0 $\rightarrow$ 0.748
686	0.019	0.051	0.051	0.7 $\rightarrow$ 0.058
698	0.279	0.623	0.514	0.85 $\rightarrow$ 0.615

We see two plausible, non-exclusive explanations for this synthetic-to-real gap, offered as hypotheses rather than established facts. First, scale of overlap: Our synthetic benchmark places 4 core communities and a comparatively small number of two-way bridge nodes, whereas the real ego-

networks contain 7-39 circles each (Table 11) - McAuley and Leskovec reported that participants identified 19 circles on average, with about 22 members each, so a typical node belongs to several circles simultaneously, in patterns considerably richer than the two-way overlap our benchmark and

our threshold rule (Eq. (6)) were designed around. A single global τ that admits a node into every community within a fixed fraction of its top membership may be too coarse an instrument for this many-small-circles regime. Second, attribute alignment: the anonymized Facebook profile features (school, employer, hometown and similar categorical fields) may correlate with only some of a user's circles (e.g. a work circle) and not others (e.g. a hobby-based circle), so the same attribute similarity S that helped separate communities cleanly on our synthetic benchmark -built so that attributes align with the community structure by construction -can inject noise into the membership score for circles the attributes do not reflect. Distinguishing between these (and other) explanations and adapting HALO's membership rule accordingly -for instance, with a per-node or per-community threshold instead of one global τ , or by learning which attributes are informative for which community -is an open problem we did not solve in this paper and flag explicitly in Section 7.

We draw a correspondingly bounded conclusion. HALO-GM is a substantial, well-validated improvement over HALO-TALB and over hard TALB [1] on attributed networks whose overlap resembles our synthetic benchmark's two-way bridge structure (Sections 6.1-6.7). On the denser, many-circle overlap structure of real social ego-networks, neither HALO-TALB nor HALO-GM currently improves on a plain hard partition and this holds even under oracle τ selection; a topology-only hard partitioner (greedy modularity) remains the strongest method we tested on this dataset. We consider this an honest characterization of HALO's current domain of validity rather than a reason to withhold the synthetic results, which remain a valid demonstration that the overlap mechanism can work well under the right conditions.

6.9 Limitations of the Study

Several limitations should be acknowledged. First, the synthetic benchmark's overlap structure (four core communities joined by two-way bridge nodes) is simpler than the overlap patterns observed in real social networks (Section 6.8), which limits the external validity of the positive synthetic results in Sections 6.1-6.7 taken alone. Second, our real-world validation uses a single dataset family (SNAP ego-Facebook circles); results may not generalize to other domains such as co-authorship or biological interaction networks, where the semantics of "attribute" and "community" differ substantially. Third, both TALB and HALO assume undirected, unweighted graphs with binary attributes; many real

networks are weighted, directed, or have continuous or textual attributes and adapting the information matrix to these settings is untested. Fourth, the overlap threshold τ and the topology-attribute weight α were tuned by grid search against ground truth on the synthetic benchmark; a fully unsupervised, data-driven selection procedure for either parameter remains an open problem and its absence limits the method's readiness for deployment on networks without labeled ground truth. Finally, the computational cost, while shown to be modest at the scales tested (up to $n \approx 1300$ nodes, Section 6.5), was not evaluated on the much larger networks (millions of nodes) common in industrial applications and the $O(n^2)$ construction of the information matrix I would require approximation or scarification to scale further.

7. CONCLUSION AND FUTURE WORK

Returning to H1 stated in Section 1: the evidence gathered here confirms that HALO's fuzzy-membership mechanism is backbone-independent and yields large accuracy gains under moderate, controlled overlap (Sections 6.1-6.7), but does not confirm that these gains hold under the denser, many-way overlap regime of real social networks (Section 6.8). H1 is therefore only partially supported: the mechanism itself generalizes across backbones as hypothesized, but its practical benefit depends on the overlap density of the target network, a boundary condition this paper identifies but does not yet resolve.

This paper introduced HALO, a fuzzy-membership mechanism that turns a hard community partition into an overlapping one by reusing an already-computed topology-attribute information matrix, with no extra similarity computation. HALO-TALB, its direct application on top of the leader-based algorithm TALB [1], improved both the Omega index and the average F1 score over hard TALB [1] on our synthetic benchmark but was matched or exceeded by a strong topology-only baseline at high noise. HALO-GM, which decouples the same mechanism from TALB [1] and re-derives it on a greedy-modularity backbone with a community-average affinity score, resolved that gap and was the strongest method tested in every synthetic noise regime. Critically, this synthetic benchmark success did not transfer to 10 real, attributed ego-Facebook networks with genuine overlapping ground truth (Section 6.8): at every threshold we tried, both HALO instantiations underperformed their own hard backbones and an oracle upper bound -selecting the best threshold per network directly against ground truth -only matched,

rather than exceeded, plain hard partitioning. We report this honestly as the central open problem rather than omitting it: on this dataset family, the limitation lies in the fuzzy-membership mechanism itself, not merely in the choice of threshold τ and we trace it to real social circles overlapping far more richly (7–39 circles per user) than the two-way bridge structure our benchmark and threshold rule were designed around.

Several directions follow directly from Section 6.8's findings. Most importantly, HALO's global, single- τ membership rule (Eq. (6)) should be replaced or augmented with a mechanism expressive enough for many-small-circles overlap -for instance a per-node adaptive threshold, a membership rule that can admit a node into several communities at genuinely different confidence levels rather than one shared cutoff, or an attribute-weighting scheme that identifies which attributes are informative for which community rather than folding all of W into a single S matrix. Second, our real-data evaluation used ten ego-networks from a single dataset [34]; a broader validation across other real, attributed, overlapping-ground-truth sources (e.g., co-authorship networks with multi-topic authors or community-labeled citation networks) would clarify whether the gap observed here is a general property of real social networks or specific to ego-Facebook circles. Third, both TALB [1] and HALO currently target undirected, unweighted networks; extending the information matrix to weighted or directed graphs is a natural next step and was also left open in the original TALB [1] paper for the non-overlapping case. Fourth, HALO-GM used greedy modularity maximisation as a representative strong backbone; repeating the same substitution with Louvain [22] or Leiden [17] is a direct extension of the recipe given in Section 4.6 and would test whether HALO-GM's synthetic-benchmark gains generalize to other strong topological partitioners and whether a stronger backbone alone -without also addressing the membership rule -is sufficient to narrow the real-data gap identified in Section 6.8.

REFERENCES:

- [1] D.-D. Lu, "Leader-based community detection algorithm in attribute networks," IEEE Access, 2021, doi: 10.1109/ACCESS.2021.3109124.
- [2] R. R. Khorasgani, J. Chen and O. R. Zaïane, "Top leaders community detection approach in information networks," in Proc. 4th SNA-KDD Workshop on Social Network Mining and Analysis, 2010.
- [3] Z. Yakoubi and R. Kanawati, "LICOD: A leader-driven algorithm for community detection in complex networks," Vietnam J. Comput. Sci., vol. 1, no. 4, pp. 241–256, 2014.
- [4] K. Shen, L. Song, X. Yang and W. Zhang, "A hierarchical diffusion algorithm for community detection in social networks," in Proc. IEEE Int. Conf. on Cyber-Enabled Distributed Computing and Knowledge Discovery, 2010, pp. 276–283.
- [5] H. Sun, H. Du, J. Huang, Y. Li, Z. Sun, L. He, X. Jia and Z. Zhao, "Leader-aware community detection in complex networks," Knowl. Inf. Syst., vol. 62, no. 2, pp. 639–668, 2020.
- [6] N. A. Helal, R. M. Ismail, N. L. Badr and M. G. M. Mostafa, "An efficient algorithm for community detection in attributed social networks," in Proc. 10th Int. Conf. on Informatics and Systems, 2016, pp. 180–184.
- [7] G. Palla, I. Derényi, I. Farkas and T. Vicsek, "Uncovering the overlapping community structure of complex networks in nature and society," Nature, vol. 435, no. 7043, pp. 814–818, 2005.
- [8] J. Xie, B. K. Szymanski and X. Liu, "SLPA: Uncovering overlapping communities in social networks via a speaker-listener interaction dynamic process," in Proc. IEEE 11th Int. Conf. on Data Mining Workshops, 2011, pp. 344–349.
- [9] J. Yang and J. Leskovec, "Overlapping community detection at scale: A nonnegative matrix factorization approach," in Proc. 6th ACM Int. Conf. on Web Search and Data Mining (WSDM), 2013, pp. 587–596.
- [10] J. J. Whang, D. F. Gleich and I. S. Dhillon, "Overlapping community detection using seed set expansion," in Proc. 22nd ACM Int. Conf. on Information and Knowledge Management (CIKM), 2013, pp. 2099–2108.
- [11] J. Xie, S. Kelley and B. K. Szymanski, "Overlapping community detection in networks: The state-of-the-art and comparative study," ACM Comput. Surv., vol. 45, no. 4, pp. 1–35, 2013.
- [12] L. M. Collins and C. W. Dent, "Omega: A general formulation of the Rand index of cluster recovery suitable for non-disjoint solutions," Multivariate Behav. Res., vol. 23, no. 2, pp. 231–242, 1988.
- [13] A. Lancichinetti, S. Fortunato and F. Radicchi, "Benchmark graphs for testing community detection algorithms," Phys. Rev. E, vol. 78, no. 4, art. 046110, 2008.

- [14] A. Lancichinetti, S. Fortunato and J. Kertész, "Detecting the overlapping and hierarchical community structure in complex networks," *New J. Phys.*, vol. 11, art. 033015, 2009.
- [15] M. Rosvall and C. T. Bergstrom, "Maps of random walks on complex networks reveal community structure," *Proc. Natl. Acad. Sci. U.S.A.*, vol. 105, no. 4, pp. 1118–1123, 2008.
- [16] M. E. J. Newman, "Modularity and community structure in networks," *Proc. Natl. Acad. Sci. U.S.A.*, vol. 103, no. 23, pp. 8577–8582, 2006.
- [17] V. A. Traag, L. Waltman and N. J. van Eck, "From Louvain to Leiden: Guaranteeing well-connected communities," *Sci. Rep.*, vol. 9, art. 5233, 2019.
- [18] A. Grover and J. Leskovec, "node2vec: Scalable feature learning for networks," in *Proc. 22nd ACM SIGKDD Int. Conf. on Knowledge Discovery and Data Mining (KDD)*, 2016, pp. 855–864.
- [19] B. Perozzi, R. Al-Rfou and S. Skiena, "DeepWalk: Online learning of social representations," in *Proc. 20th ACM SIGKDD Int. Conf. on Knowledge Discovery and Data Mining (KDD)*, 2014, pp. 701–710.
- [20] J. Yang, J. McAuley and J. Leskovec, "Community detection in networks with node attributes," in *Proc. IEEE 13th Int. Conf. on Data Mining (ICDM)*, 2013, pp. 1151–1156.
- [21] C. Wang, S. Pan, R. Hu, G. Long, J. Jiang and C. Zhang, "Attributed graph clustering: A deep attentional embedding approach," in *Proc. 28th Int. Joint Conf. on Artificial Intelligence (IJCAI)*, 2019, pp. 3670–3676.
- [22] V. D. Blondel, J.-L. Guillaume, R. Lambiotte and E. Lefebvre, "Fast unfolding of communities in large networks," *J. Stat. Mech.*, vol. 2008, art. P10008, 2008.
- [23] M. E. J. Newman, "The structure and function of complex networks," *SIAM Rev.*, vol. 45, no. 2, pp. 167–256, 2003.
- [24] Y. Zhou, H. Cheng and J. X. Yu, "Graph clustering based on structural/attribute similarities," *Proc. VLDB Endowment*, vol. 2, no. 1, pp. 718–729, 2009.
- [25] Z. Xu, Y. Ke, Y. Wang, H. Cheng and J. Cheng, "A model-based approach to attributed graph clustering," in *Proc. 2012 ACM SIGMOD Int. Conf. on Management of Data*, 2012, pp. 505–516.
- [26] K. Steinhaeuser and N. V. Chawla, "Identifying and evaluating community structure in complex networks," *Pattern Recognit. Lett.*, vol. 31, no. 5, pp. 413–421, 2010.
- [27] S. Ahajjam, M. El Haddad and H. Badir, "A new scalable leader-community detection approach for community detection in social networks," *Social Networks*, vol. 54, pp. 41–49, 2018.
- [28] N. A. Helal, R. M. Ismail, N. L. Badr and M. G. M. Mostafa, "Leader-based community detection algorithm for social networks," *WIRES Data Mining Knowl. Discov.*, vol. 7, no. 6, 2017.
- [29] A. Goyal, F. Bonchi and L. V. S. Lakshmanan, "Discovering leaders from community actions," in *Proc. 17th ACM Conf. on Information and Knowledge Management (CIKM)*, 2008, pp. 499–508.
- [30] Y. Li, C. Sha, X. Huang and Y. Zhang, "Community detection in attributed graphs: An embedding approach," in *Proc. 32nd AAAI Conf. on Artificial Intelligence*, 2018, pp. 338–345.
- [31] A. Clauset, M. E. J. Newman and C. Moore, "Finding community structure in very large networks," *Phys. Rev. E*, vol. 70, no. 6, art. 066111, 2004.
- [32] M. Ester, H.-P. Kriegel, J. Sander and X. Xu, "A density-based algorithm for discovering clusters in large spatial databases with noise," in *Proc. 2nd Int. Conf. on Knowledge Discovery and Data Mining (KDD)*, 1996, pp. 226–231.
- [33] D. Huang, J.-H. Lai and C.-D. Wang, "Robust ensemble clustering using probability trajectories," *IEEE Trans. Knowl. Data Eng.*, vol. 28, no. 5, pp. 1312–1326, 2016.
- [34] J. McAuley and J. Leskovec, "Learning to discover social circles in ego networks," in *Proc. 26th Int. Conf. on Neural Information Processing Systems (NeurIPS)*, 2012, pp. 539–547.