

Exploring multi-layered networks through random walks: bridging offline optimization and online learning

Xiangxiang Dai ^a, Xutong Liu ^{b,*}, Jinhang Zuo ^c, Xiaowei Chen ^d, Wei Chen ^e, John C S Lui ^a

^a Department of Computer Science and Engineering, The Chinese University of Hong Kong, Hong Kong SAR, China

^b Department of Electrical and Computer Engineering, University of Washington, Tacoma, USA

^c Department of Computer Science, City University of Hong Kong, Hong Kong SAR, China

^d Bytedance, USA

^e Microsoft Research Asia, China

ARTICLE INFO

Keywords:

Network exploration
Random walk
Offline optimization
Online learning
Combinatorial multi-armed bandit

ABSTRACT

The multi-layered network exploration problem MuLaNE is a significant challenge derived from various practical applications. In MuLaNE, there are multiple layers of the network, where each node is assigned an importance weight, and each layer is explored through random walks. The objective is to allocate a total random walk budget B across the network layers to maximize the cumulative weights of the unique nodes visited. We systematically approach this problem by addressing both offline optimization and online learning scenarios. For the offline optimization case, where the network structure and node weights are known, we propose constant-ratio approximation algorithms based on greedy strategies for overlapping networks, and exact optimal solutions using greedy or dynamic programming for non-overlapping networks. In the online learning scenario, where neither the network structure nor the node weights are initially known, we adapt the combinatorial multi-armed bandit framework. We develop algorithms with two distinct confidence radius designs to simultaneously learn random walk parameters and node weights, while optimizing budget allocation across multiple rounds, achieving logarithmic regret bounds. Finally, experimental results on real-world social and computer networks validate the practical applicability of MuLaNE and our theoretical findings.

1. Introduction

Network exploration is a fundamental process for discovering information and resources located at nodes within a network, where random walks are frequently employed as an effective method [1–3]. In this paper, we investigate the problem of multi-layered network exploration via random walks, a model applicable to various real-world scenarios such as resource discovery in peer-to-peer (P2P) networks and web navigation in online social networks (OSNs), as depicted in Fig. 1.

In P2P networks, as shown in Fig. 1a, a user often needs to find valuable resources (like files or data) that are scattered across many different peers, which may belong to several distinct networks. A common and effective decentralized search strategy is to launch multiple parallel random walkers [1,4], where each walker is a query that propagates from peer to peer. Critically, these

* Corresponding author.

E-mail addresses: xxdai23@cse.cuhk.edu.hk (X. Dai), xutongl@uw.edu (X. Liu), jinhang.zuo@cityu.edu.hk (J. Zuo), weic@microsoft.com (W. Chen).

<https://doi.org/10.1016/j.artint.2026.104500>

Received 28 November 2024; Received in revised form 11 September 2025; Accepted 12 February 2026

Available online 26 February 2026

0004-3702/© 2026 Published by Elsevier B.V.

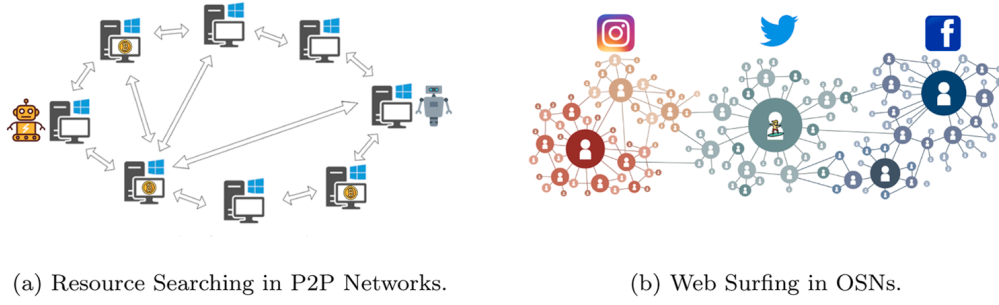


Fig. 1. Applications on network exploration via random walks.

search queries don't run indefinitely; they operate under strict constraints, such as a maximum hop count or a time-to-live limit. Given that resources vary in importance and are distributed unevenly, the central challenge becomes one of resource allocation: *how should a user with a limited total search capacity strategically divide this effort among different walkers in various networks to maximize the total value of the unique, important resources they discover?*

A similar challenge emerges in the automated exploration of online social networks (OSNs), where autonomous explorers navigate social connections to uncover information [5,6]. Unlike tasks with predefined targets, these explorers lack a feasible list of nodes due to OSNs' vast scale and dynamism; their core objective is discovering unknown content. Instead, they traverse platforms like Instagram, Twitter, and Facebook through real-time programmatic navigation, guided by emerging connections (Fig. 1b). This path-dependent process—starting at a seed node, following edges (e.g., likes or connections) to profiles, and continuing onward—naturally models as a random walk. Yet, it is limited by a finite computational budget, stemming from practical constraints like API rate limits and per-step processing costs. This poses a central optimization question: *How can an automated system optimally allocate its budget across platforms to maximize high-value information discovery?*

We formalize the above application scenarios as the **Multi-Layered Network Exploration (MuLaNE)** problem. In MuLaNE, the overall network to be explored (e.g., a combination of different OSNs) is modeled as a multi-layered network $\mathcal{G}$ with m layers $L_1, \dots, L_m$. Each layer L_i (e.g., an individual OSN) is a weighted directed graph $G_i = (\mathcal{V}_i, \mathcal{E}_i)$, where $\mathcal{V}_i$ is the set of nodes (e.g., user pages with importance weights σ_u), and $\mathcal{E}_i$ is the set of directed edges with associated edge weight function w_i . Each layer is explored by an explorer (or random walker) W_i , which starts according to a given distribution α_i and follows the outgoing edges with probabilities proportional to their edge weights. Each step consumes one unit of budget¹. Given a total budget B and per-layer limits c_i , the goal is to determine the optimal budget allocation $\mathbf{k} = (k_1, \dots, k_m)$ within the budget constraints to maximize the expected total weight of *unique* visited nodes.

In practical applications, the network structure $\mathcal{G}$, the node weights σ , and the starting distributions α_i may not be known in advance. Therefore, we consider both the *offline optimization* scenario, where $\mathcal{G}$, σ , and α_i are known, and the *online learning* scenario, where these elements are unknown. Furthermore, different layers may have overlapping vertices (e.g., the same user's Homepage may appear in multiple OSNs), and the starting distributions α_i may or may not correspond to stationary distributions. This paper provides a systematic study of all these combinations of cases:

- For the offline optimization process, we first address the general case with overlapping layers and propose constant approximation algorithms. These are based on the lattice submodularity property for non-stationary starting distributions α_i , and the diminishing returns (DR) submodularity property for stationary α_i . For the special case of non-overlapping layers, we design a dynamic programming algorithm that computes the exact optimal solution for MuLaNE.
- In the online learning setting, exploration is conducted over multiple rounds, each constrained by the same budget B and individual layer budgets c_i . After each round, the reward is the total weight of the unique nodes visited, and the trajectory of each explorer, along with the importance weights of the visited nodes, is observed. This feedback informs future exploration rounds by providing valuable insights into the network's structure.

We adapt the combinatorial multi-armed bandit (CMAB) framework [7] to minimize regret, defined as the difference between the cumulative reward achieved by “the exact or approximate offline algorithm” and that achieved by “the online learning algorithm” over T rounds. For the offline optimization setting, the core challenge is to solve a complex combinatorial optimization problem for budget allocation, given full knowledge of the network. The problem is NP-hard in general, as the reward function may not possess desirable properties. To address this, we leverage the property of lattice submodularity to design **constant-ratio approximation algorithms** for the general case, and we provide **the exact optimal solutions** via greedy or dynamic programming methods for important special cases. For the online learning setting, the network structure and node weights are initially unknown, and the challenge lies in sequentially allocating the budget over multiple rounds. This requires balancing the exploration of unknown parameters with the exploitation of current knowledge, with the goal of minimizing cumulative regret against the offline solution. To tackle this, we develop **novel online algorithms** that learn random walk parameters, i.e., visiting probabilities, and node weights simultaneously to achieve **logarithmic regret bounds**.

¹ The model can be extended to allow multiple steps per budget unit.

Contributions. In summary, our contributions are as follows: (a) We model the multi-layered network exploration via the random walks problem (MuLaNE) as an abstraction for various real-world applications. (b) We conduct a systematic study of the MuLaNE problem, designing specialized algorithms and providing rigorous theoretical analysis for both offline and online settings, considering overlapping and non-overlapping multi-layered networks. (c) We perform empirical validation on real-world social and computer network datasets, demonstrating the practical applicability of MuLaNE and the effectiveness of our proposed algorithms.

Paper Organization. Section 3 presents the formal problem setting of MuLaNE. In Section 4, we introduce the equivalent bipartite coverage model to derive the explicit form of the reward function. Section 5 discusses offline algorithms for four distinct scenarios, providing provable approximation guarantees and running-time analyses. In Section 6, we present different online learning algorithms along with regret analysis for both overlapping and non-overlapping cases. Empirical results are provided in Section 7, followed by related work in Section 2. Finally, Section 8 concludes the paper.

2. Related work

Network exploration via random walks has been studied in various application contexts such as centrality measuring [2], large-scale network sampling [8], and influence maximization [3]. Multiple random walks has also been used, e.g. in Lv et al. [1] for query resolution in peer-to-peer networks. However, none of these studies addresses the budgeted MuLaNE problem. MuLaNE is also related to the influence maximization (IM) problem [9,10] and can be viewed as a special variant of IM. Different from the standard Independent Cascade (IC) model [11] in IM, which randomly broadcasts to all its neighbors, the propagation process of MuLaNE is a random walk that selects one neighbor. Moreover, each unit of budget in MuLaNE is used to propagate one random-walk step, not to select one seed node as in IC.

While the MuLaNE problem formulation is novel, its offline version can be framed as an instance of a budget allocation problem. Our work, therefore, builds upon the theoretical tools developed for this class of problems. The offline budget allocation problem has been studied by Alon et al. [12], Soma et al. [13], the latter of which propose the lattice submodularity and a constant approximation algorithm based on this property. Our offline overlapping MuLaNE setting is based on a similar approach, but we provide a better approximation ratio compared to the original analysis in Alon et al. [12]. Moreover, we further show that the stationary setting enjoys DR-submodularity, leading to an efficient algorithm with a better approximation ratio.

The multi-armed bandit (MAB) problem is first studied by Robbins [14] and then extended by many studies (cf. [15]). Our online setting fits into the general Combinatorial MAB (CMAB) framework of Chen et al. [7], Gai et al. [16]. CUCB is proposed as a general algorithm for CMAB [7,17]. Our study includes several adaptations, such as handling unknown node weights, defining intermediate random variables as base arms, and so on. The community exploration problem, as addressed by Chen et al. [18], represents a special case of MuLaNE that is restricted to non-overlapping complete graphs. Consequently, their approach is inherently different from and simpler than the generalized methods we propose. While Chen et al. [7], Wang and Chen [17], Liu et al. [19] can also tackle the MuLaNE problem, their algorithms' regret upper bound scales linearly with the number of explorers m . Our explorer-relaxed online algorithm improves upon this result, achieving a tighter bound by a factor of $\tilde{O}\left(\frac{m}{\log m}\right)$.

3. Preliminaries

Basic Notations. In this paper, we use $\mathbb{R}$ and $\mathbb{Z}$ to denote the sets of real numbers and integers, respectively, with subscripts ≥ 0 and > 0 representing their non-negative and positive subsets. We consider a multi-layered network $\mathcal{G}(\mathcal{V}, \mathcal{E}, \sigma)$, consisting of m layers $L_1, \dots, L_m$ and N unique nodes in the set $\mathcal{V}$, with fixed node weights $\sigma \in [0, 1]^{\mathcal{V}}$. Let $[m]$ denote the set $1, 2, \dots, m$. Each layer L_i , for $i \in [m]$, is a subgraph of $\mathcal{G}$ and is represented as a weighted directed graph $\mathcal{G}_i(\mathcal{V}_i, \mathcal{E}_i; w_i)$, where $\mathcal{V}_i \subseteq \mathcal{V}$ is the set of vertices in layer L_i , $\mathcal{E}_i \subseteq \mathcal{V}_i \times \mathcal{V}_i$ is the set of edges, and $w_i : \mathcal{E}_i \rightarrow \mathbb{R}_{>0}$ is the edge weight function for the edges in $\mathcal{G}_i$. Two layers L_i and L_j are considered *non-overlapping* if $\mathcal{V}_i \cap \mathcal{V}_j = \emptyset$. The network $\mathcal{G}$ is called a *non-overlapping* multi-layered network if all layers are non-overlapping; otherwise, it is an *overlapping* multi-layered network. Fig. 2 illustrates both types of networks. Let $n_i = |\mathcal{V}_i|$ denote the size of layer L_i , and the adjacency matrix of $\mathcal{G}_i$ is defined as $\mathbf{A}_i \in \mathbb{R}_{\geq 0}^{n_i \times n_i}$, where $\mathbf{A}_i[u, v] = w_i((u, v))$ if $(u, v) \in \mathcal{E}_i$ and 0 otherwise. Based on the definition of $\mathbf{A}_i$, we also define the transition probability matrix $\mathbf{P}_i \in \mathbb{R}_{\geq 0}^{n_i \times n_i}$ as $\mathbf{P}_i[u, v] := \mathbf{A}_i[u, v] / (\sum_{w: (u, w) \in \mathcal{E}_i} \mathbf{A}_i[u, w])$.

Exploration Rule. Each layer L_i is associated with an explorer W_i , a budget $k_i \in \mathbb{Z}_{\geq 0}$ and an initial starting distribution $\alpha_i = (\alpha_{i,u})_{u \in \mathcal{V}_i} \in \mathbb{R}_{\geq 0}^{n_i}$ with $\sum_{u \in \mathcal{V}_i} \alpha_{i,u} = 1$. The exploration process of W_i is identical to applying *weighted random walks* within layer L_i . Specifically, the explorer W_i starts the exploration from a random node $v_1 \in \mathcal{V}_i$ with probability α_{i,v_1} ; it consumes one unit of budget and continues the exploration by walking from v_1 to one of its out-neighbors, say v_2 , with the transition probability matrix $\mathbf{P}_i[v_1, v_2]$. Consider visiting the initial node v_1 as the first step, the process is repeated until the explorer W_i walks k_i steps and visits k_i nodes (with possibly duplicated nodes). Note that our results can be easily generalized by allowing W_i to take $\lambda_i \in \mathbb{Z}_{>0}$ steps per unit of budget. It can be readily verified that this modification does not affect the algorithm or its analysis. For simplicity, we set $\lambda_i = 1$.

Reward Function. We define the exploration trajectory for explorer W_i after exploring k_i nodes as $\Phi(i, k_i) := (X_{i,1}, \dots, X_{i,k_i})$, where $X_{i,j} \in \mathcal{V}_i$ denotes the node visited at j -th step, and $\Pr(X_{i,1}=u) = \alpha_{i,u}, u \in \mathcal{V}_i$. Note that $\alpha_{i,u}$ denotes the probability of node u being the initial starting node for W_i . The reward for $\Phi(i, k_i)$ is defined as the total weights of *unique* nodes visited by W_i , i.e., $\sum_{v \in \bigcup_{j=1}^{k_i} \{X_{i,j}\}} \sigma_v$. Considering trajectories of all W_i 's, the total reward is the total weights of *unique* nodes visited by all random walkers, i.e., $\sum_{v \in \bigcup_{i=1}^m \bigcup_{j=1}^{k_i} \{X_{i,j}\}} \sigma_v$. For notational simplicity, let $\mathcal{G} := (\mathcal{G}_1, \dots, \mathcal{G}_m)$, $\alpha := (\alpha_1, \dots, \alpha_m)$ and $\sigma = (\sigma_1, \dots, \sigma_{|\mathcal{V}|})$ be the parameters of a

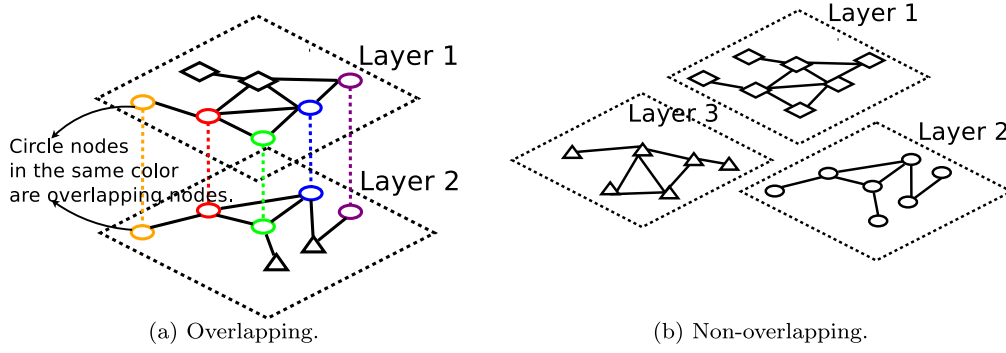


Fig. 2. Two types of multi-layered networks.

problem instance and $\mathbf{k} := (k_1, \dots, k_m)$ be the allocation vector. For overlapping multi-layered network $\mathcal{G}$, the total expected reward is

$$r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k}) := \mathbb{E}_{\Phi(1, k_1), \dots, \Phi(m, k_m)} \left[\sum_{v \in \bigcup_{i=1}^m \bigcup_{j=1}^{k_i} \{X_{i,j}\}} \sigma_v \right], \quad (1)$$

where its explicit formula will be discussed later in Section 4. For non-overlapping $\mathcal{G}$, there are no overlapping nodes among different layers, i.e., a node can only be visited by the random explorer from the same layer. Hence, the reward function can be simplified and written as the summation over separated layers,

$$r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k}) := \sum_{i=1}^m \mathbb{E}_{\Phi(i, k_i)} \left[\sum_{v \in \bigcup_{j=1}^{k_i} \{X_{i,j}\}} \sigma_v \right]. \quad (2)$$

Problem Formulation. We are interested in the *budget allocation problem*: how to distribute a total budget $B \in \mathbb{Z}_{\geq 0}$ among m explorers to maximize the total weight of unique nodes visited. Additionally, each explorer $i \in [m]$ has an individual budget constraint c_i , meaning the allocated budget k_i must satisfy $k_i \leq c_i$. We assume $c_i \leq B$ and $\sum_{i=1}^m c_i \geq B$, as the problem becomes trivial otherwise. Specifically, if $\sum_{i=1}^m c_i < B$, the optimal solution is simply to allocate $k_i = c_i$ for all explorer i , which satisfies the total budget constraint.

Definition 1. Given graph structures $\mathcal{G} := (\mathcal{G}_1, \dots, \mathcal{G}_m)$, total budget B , starting distributions $\alpha := (\alpha_1, \dots, \alpha_m)$, and budget constraints $\mathbf{c} := (c_1, \dots, c_m)$, the **Multi-Layered Network Exploration problem**, denoted as MuLaNE, is formulated as the following optimization problem,

$$\max_{\mathbf{k}} r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k}) \text{ s.t. } \mathbf{k} \in \mathbb{Z}_{\geq 0}^m \leq \mathbf{c}, \sum_{i=1}^m k_i \leq B, \quad (3)$$

where $r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k})$ is given by Eq. (1) or Eq. (2).

We consider the following two settings for the MuLaNE problem:

Offline Setting. For the offline setting, all the parameters of the problem instance $(\mathcal{G}, \alpha, \sigma, \mathbf{c}, B)$ are given in advance. The objective is to determine the optimal allocation of the budget $\mathbf{k}^*$ by solving the optimization problem defined in Eq. (3).

Online Setting. For the online setting, we consider T -round explorations. Prior to exploration, only an upper bound on the total number of nodes in $\mathcal{G}$ and the number of layers m are known². The network structure $\mathcal{G}$, starting distributions α , and node weights σ are initially unknown. In round $t \in [T]$, we choose the budget allocation $\mathbf{k}_t := (k_{t,1}, \dots, k_{t,m})$ based solely on observations from previous rounds, where $k_{t,i} \in \mathbb{Z}_{\geq 0}$ is the budget allocated to the i -th layer and $\sum_{i=1}^m k_{t,i} \leq B, \mathbf{k}_t \leq \mathbf{c}$. By choosing *action* $\mathbf{k}_t$, each explorer W_i performs $k_{t,i}$ exploration steps on layer L_i , generating an exploration trajectory $\Phi(i, k_{t,i}) = (X_{i,1}, \dots, X_{i,k_{t,i}})$. After the action, the exploration trajectories $\Phi(i, k_{t,i})$ and the fixed importance weights σ_u for all nodes $u \in \Phi(i, k_{t,i})$ are revealed as *feedback* (Random node weights with unknown mean vector σ is considered in Appendices F and G). Feedback results are used to learn parameters related to the graph structures $\mathcal{G}$, the starting distribution α , and the node weights σ . This information guides the selection of better budget allocations in subsequent rounds. In round t , the reward obtained is the total weights of unique nodes visited by all random explorers. Our objective is to design an efficient online learning algorithm, denoted as A , to guide the budget allocation process and maximize the cumulative reward over T rounds.

A key challenge for the online algorithm is managing the exploration-exploitation tradeoff. The performance of an online learning algorithm is commonly measured by cumulative regret. Formally, the T -round (ξ, β) -approximation regret of algorithm A is defined as:

$$\text{Reg}_{\mathcal{G}, \alpha, \sigma}(T) = \xi \beta T \cdot r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k}^*) - \mathbb{E} \left[\sum_{t=1}^T r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k}_t^A) \right], \quad (4)$$

where $r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k}^*)$ is the reward value for the optimal budget allocation $\mathbf{k}^*$, $\mathbf{k}_t^A$ is budget allocation selected by the learning algorithm A in round t , and the expectation is taken over the randomness of both the algorithm and the exploration trajectories across all T rounds. The constants (ξ, β) represent the approximation guarantee of an offline oracle, as explained below. Similar to other online learning

² More precisely we only need to know the number of independent explorers. If multiple explorers operate on the same layer, this is equivalent to treating them as separate layers with identical graph structures.

frameworks [7,17,20], we assume that the online learning algorithm has access to an offline (ξ, β) -approximation oracle, which for problem instance $(\mathcal{G}, \alpha, \sigma, c, B)$ outputs an action k such that $\Pr(r_{\mathcal{G}, \alpha, \sigma}(k) \geq \xi \cdot r_{\mathcal{G}, \alpha, \sigma}(k^*)) \geq \beta$. Importantly, in our actual online learning algorithms, we do not try to recover the graph structure $\mathcal{G}$ and the starting distribution α but take certain intermediate parameters as inputs.

Submodularity and DR-Submodularity Over Integer Lattices. To address the MuLaNE problem, we exploit the submodular and diminishing returns (DR)-submodular properties of the reward function. Let $\mathbb{R}$ and $\mathbb{Z}_{\geq 0}$ denote the set of reals and non-negative integers. For any $\mathbf{x}, \mathbf{y} \in \mathbb{Z}_{\geq 0}^m$, we denote $\mathbf{x} \vee \mathbf{y}, \mathbf{x} \wedge \mathbf{y} \in \mathbb{Z}_{\geq 0}^m$ as the coordinate-wise maximum and minimum of these two vectors, i.e., $(\mathbf{x} \vee \mathbf{y})_i = \max\{x_i, y_i\}, (\mathbf{x} \wedge \mathbf{y})_i = \min\{x_i, y_i\}$. We define a function $f : \mathbb{Z}_{\geq 0}^m \rightarrow \mathbb{R}$ over the integer lattice $\mathbb{Z}_{\geq 0}^m$ as a *submodular* function if the following inequality holds for any $\mathbf{x}, \mathbf{y} \in \mathbb{Z}_{\geq 0}^m$:

$$f(\mathbf{x} \vee \mathbf{y}) + f(\mathbf{x} \wedge \mathbf{y}) \leq f(\mathbf{x}) + f(\mathbf{y}). \quad (5)$$

Let $\text{supp}^+(\mathbf{x} - \mathbf{y})$ denote the set $I = \{i \in [m] : x_i > y_i\}$, $\chi_I = (0, \dots, 1, \dots, 0)$ be the one-hot vector whose i -th element is 1 and 0 otherwise, and $\mathbf{x} \geq \mathbf{y}$ means $x_i \geq y_i$ for all $i \in [m]$. We define a function $f : \mathbb{Z}_{\geq 0}^m \rightarrow \mathbb{R}$ as a *DR-submodular* function if the following inequality holds for any $\mathbf{x} \leq \mathbf{y}$ and $i \in [m]$,

$$f(\mathbf{y} + \chi_i) - f(\mathbf{y}) \leq f(\mathbf{x} + \chi_i) - f(\mathbf{x}). \quad (6)$$

We say a function f is *monotone* if for any $\mathbf{x} \leq \mathbf{y}$, $f(\mathbf{x}) \leq f(\mathbf{y})$. Without loss of generality, we assume monotone functions are normalized, i.e., $f(\mathbf{0}) = 0$. For a function $f : \{0, 1\}^m \rightarrow \mathbb{R}$ defined on a set, submodularity and DR-submodularity are equivalent. On general integer lattices, however, the distinction is critical. A DR-submodular function $f : \mathbb{Z}_{\geq 0}^m \rightarrow \mathbb{R}$ is always a submodular function, but the converse is not necessarily true. Specifically, DR-submodularity implies that the marginal gain from a single-unit increase in any coordinate is non-increasing. The more general lattice submodularity does not require this strict, per-unit diminishing return property.

4. Equivalent bipartite coverage model

In this section, we derive the formulation of the reward function, a crucial step for both offline optimization and online learning. We first introduce an equivalent bipartite coverage model to formalize the problem, then use this framework to express the reward function mathematically, and finally analyze the properties of its core component: the node visiting probability.

4.1. Formulation of reward function

To derive the explicit formulation of the reward function $r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k})$ in Eq. (1), we construct an undirected bipartite coverage graph $B(\mathcal{W}, \mathcal{V}, \mathcal{E}')$, where $\mathcal{W} = \{W_1, \dots, W_m\}$ denotes m random explorers, $\mathcal{V} = \bigcup_{i \in [m]} \mathcal{V}_i$ denotes all possible distinct nodes to be explored in $\mathcal{G}$, and the edge set $\mathcal{E}' = \{(W_i, u) | u \in \mathcal{V}_i, i \in [m]\}$ indicating whether node u could be visited by W_i . For each edge $(W_i, u) \in \mathcal{E}'$, we associate it with $c_i + 1$ visiting probabilities denoted as $P_{i,u}(k_i)$ for $k_i \in \{0\} \cup [c_i]$. $P_{i,u}(k_i)$ represents the probability that the node u is visited by the random walker W_i given the budget k_i , i.e., $P_{i,u}(k_i) = \Pr(u \in \Phi(i, k_i))$. We will provide the analytical formula and detailed properties of these visiting probabilities later on. Since W_i can never visit the node outside the i -th layer (i.e., $u \notin \mathcal{G}_i$), we set $P_{i,u}(k_i) = 0$ if $(W_i, u) \notin \mathcal{E}'$. Then, given budget allocation $\mathbf{k}$, the probability of a node u visited by at least one random explorer is $\Pr(u, \mathbf{k}) = 1 - \prod_{i \in [m]} (1 - P_{i,u}(k_i))$ because each W_i has the independent probability $P_{i,u}(k_i)$ to visit u . By summing over all possible nodes, the reward function is

$$r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k}) = \sum_{u \in \mathcal{V}} \sigma_u \left(1 - \prod_{i \in [m]} (1 - P_{i,u}(k_i)) \right) \quad (7)$$

According to Eq. (2), for a non-overlapping multi-layered network, we can rewrite the reward function as:

$$r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k}) = \sum_{i \in [m]} \sum_{u \in \mathcal{V}_i} \sigma_u P_{i,u}(k_i). \quad (8)$$

Remark. The bipartite coverage model abstracts the integration of graph structure and random-walk exploration into $P_{i,u}(k_i)$, thereby determining the reward function. Its generality allows seamless extensions, e.g., adding a new random walker to a layer corresponds to inserting a node in $\mathcal{W}$ and associated edges in $\mathcal{E}'$, without altering the fundamental properties of the reward function or the applicability of our optimization and learning techniques.

4.2. Properties of the visiting probability

The quantities $P_{i,u}(k_i)$'s in Eqs. (7) and (8) are crucial for both our offline and online algorithms, and thus we provide their analytical formulas and properties here. To analyze the property of the reward function, we apply the absorbing Markov Chain technique to calculate $P_{i,u}(k_i)$. For $u \notin \mathcal{G}_i$, $P_{i,u}(k_i) = 0$; For $u \in \mathcal{G}_i$, we create an absorbing Markov Chain $\mathbf{P}_i(\mathbf{u}) \in \mathbb{R}^{n_i \times n_i}$ by setting the target node u as the absorbing node. We derive the corresponding transition matrix $\mathbf{P}_i(\mathbf{u})$ as $\mathbf{P}_i(\mathbf{u})[v, \cdot] = \chi_u^\top$ if $v = u$ and $\mathbf{P}_i(\mathbf{u})[v, \cdot] = \mathbf{P}_i[v, \cdot]$ otherwise, where $\mathbf{P}_i[v, \cdot]$ denotes the row vector corresponding to the node v of $\mathbf{P}_i$, and $\chi_u = (0, \dots, 0, 1, 0, \dots, 0)^\top$ denotes the one-hot vector with 1 at the u -th entry and 0 elsewhere.³ Intuitively, $\mathbf{P}_i(\mathbf{u})$ corresponds to the transition probability matrix of $\mathcal{G}_i$ after removing all out-edges of u and adding a self-loop to itself in the original graph $\mathcal{G}_i$. The random walker W_i will be trapped in u if

³ When related to matrix operations, we treat all vectors as column vectors by default.

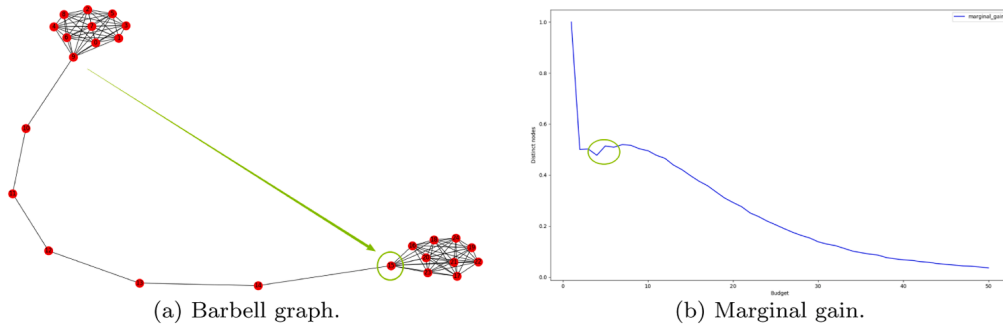


Fig. 3. An example showing $g_{i,u}(k_i)$ is not monotone.

it ever visits u . We observe that $P_{i,u}(k_i)$ equals to the probability the random walker stays in the absorbing node u at step $k_i \in \mathbb{Z}_{>0}$ (trivially, $P_{i,u}(0) = 0$), if it starts from the initial distribution α_i , i.e.,

$$P_{i,u}(k_i) = \alpha_i^\top P_i(u)^{k_i-1} \chi_u. \quad (9)$$

To see this, the probability that node u has not been visited before (included) step k equals to $\alpha P(u)^{k-1}(\mathbf{1} - \chi_u)$, where $\mathbf{1} \in \mathbb{R}^{1 \times n}$ is the all 1 vector. From this, we have $1 - \Pr(u, k) = 1 - \alpha P(u)^{k-1} \chi_u$. Then define the marginal gain of $P_{i,u}(\cdot)$ at step $k_i \geq 1$ as,

$$g_{i,u}(k_i) = P_{i,u}(k_i) - P_{i,u}(k_i - 1). \quad (10)$$

The physical meaning of $g_{i,u}(k_i)$ is the probability that node u is visited exactly at the k_i -th step and not visited before k_i . Now, we can show $g_{i,u}(k_i)$ is non-negative,

Lemma 1. $g_{i,u}(k_i) \geq 0$ for any $i \in [m]$, $u \in \mathcal{V}_i$, $k_i \in \mathbb{Z}_{>0}$.

This means $P_{i,u}(k_i)$ is non-decreasing with respect to step k_i . In other words, the more budgets we allocate to W_i , the higher probability W_i can visit u in k_i steps.

Starting from Arbitrary Distributions: Although $P_{i,u}(k_i)$ is monotone, i.e., non-decreasing, w.r.t k_i , $g_{i,u}(k_i)$ may not be monotonic non-increasing under arbitrary staring distributions. Fig. 3 illustrates this using a classic barbell graph, shown in Fig. 3a, with two densely connected clusters joined by a single path. A random walk starting at one end is likely confined to the first cluster initially, yielding negligible marginal gain for reaching a target node in the distant cluster. However, once the walk crosses the connecting path—a low-probability event requiring sufficient budgets—it gains access to the entire second cluster, causing a sharp spike in marginal gain. This non-monotonic behavior is depicted in Fig. 3b, which shows that $g_{i,u}(k_i)$ is not always non-increasing. Intuitively, there may exist one critical step k_i such that $P_{i,u}(k_i)$ suddenly increases by a large value (see Section B.1). This means that $P_{i,u}(k_i)$ lacks the diminishing return (or “discretely concave”) property, which many problems rely on to provide good optimization [21,22]. In other words, we are dealing with a more challenging non-concave optimization problem for the discrete budget allocation.

Starting from Stationary Distributions: If our walkers start with the stationary distribution, the following property holds: $g_{i,u}(k_i)$ is non-increasing w.r.t k_i .

Lemma 2. $g_{i,u}(k_i + 1) - g_{i,u}(k_i) \leq 0$ for any $i \in [m]$, $u \in \mathcal{V}_i$, $k_i \in \mathbb{Z}_{>0}$, if $\alpha_i = \pi_i$, where $\pi_i^\top P_i = \pi_i^\top$.

More interestingly, the stationary distribution π_i is the only starting distribution for P_i such that for any $u \in \mathcal{V}_i$, $k_i \in \mathbb{Z}_{>0}$, $g_{i,u}(k_i + 1) - g_{i,u}(k_i) \leq 0$ when P_i is ergodic (the random walker can visit any node in finite steps starting from any node) and there are no self loops (the random walker will not stay in the same node in any two consecutive steps) in $\mathcal{G}_i$. This means that the diminishing return property is missing for most starting distributions. The proofs are included in the Appendix B.

5. Offline optimization for MuLaNE

This section develops our optimization algorithms for the offline MuLaNE problem, which serve as the foundation and performance benchmark for the online learning setting. Our approach is to systematically analyze different scenarios, starting with the most general case and then moving to more specialized settings where stronger theoretical guarantees can be achieved. First, we consider the general case, where layers are overlapping and starting distributions are arbitrary. Next, we consider the starting distribution as the stationary distribution, and give solutions with better solution quality and time complexity for the approximation algorithms. Then we analyze special cases where layers are non-overlapping and give exact optimal solutions for arbitrary distributions. Table 1 summarizes the offline models and algorithmic results.

5.1. Offline algorithm for overlapping mulane

Starting from Arbitrary Distributions. Based on the equivalent bipartite coverage model, one can observe that our problem formulation is a generalization of the Probabilistic Maximum Coverage (PMC) [7,23,24] problem, which is NP-hard and has a $(1 - 1/e)$

Table 1
Summary of the offline models and algorithms.

Overlapping?	Starting distribution	Algorithm	Apprx ratio	Time complexity
✓	Arbitrary	Budget Effective Greedy	$(1 - e^{-\eta})^a$	$O(B\ c\ _{\infty}mn_{\max} + \ c\ _{\infty}mn_{\max}^3)$
✓	Stationary	Myopic Greedy	$(1 - 1/e)$	$O(Bmn_{\max} + \ c\ _{\infty}mn_{\max}^3)$
×	Arbitrary	Dynamic Programming	1	$O(B\ c\ _{\infty}m + \ c\ _{\infty}mn_{\max}^3)$
×	Stationary	Myopic Greedy	1	$O(B \log m + \ c\ _{\infty}mn_{\max}^3)$

^a $1 - e^{-\eta} \approx 0.357$, where η is the solution of $e^{\eta} = 2 - \eta$.

Algorithm 1 Budget effective greedy (BEG) algorithm for the overlapping MuLaNE.

Input: Network $\mathcal{G}$, starting distributions α , node weights σ , budget B , constraints c .

Output: Budget allocation k .

1: Compute visiting probabilities $(P_{i,u}(b))_{i \in [m], u \in \mathcal{V}, b \in [c_i]}$ according to Eq. (9).

2: $k \leftarrow \text{BEG}((P_{i,u}(b))_{i \in [m], u \in \mathcal{V}, b \in [c_i]}, \sigma, B, c)$.

3: **Procedure** $\text{BEG}((P_{i,u}(b))_{i \in [m], u \in \mathcal{V}, b \in [c_i]}, \sigma, B, c)$

4: Let $k := (k_1, \dots, k_m) \leftarrow \mathbf{0}$, $K \leftarrow B$.

5: Let $Q \leftarrow \{(i, b_i) \mid i \in [m], 1 \leq b_i \leq c_i\}$.

6: **while** $K > 0$ and $Q \neq \emptyset$ **do**

7: $(i^*, b^*) \leftarrow \arg \max_{(i,b) \in Q} \delta(i, b, k)/b$. {Eq. (11)}

8: $k_{i^*} \leftarrow k_{i^*} + b^*$, $K \leftarrow K - b^*$.

9: Modify all pairs $(i, b) \in Q$ to $(i, b - b^*)$.

10: Remove all pairs $(i, b) \in Q$ such that $b \leq 0$.

11: **end while**

12: **for** $i \in [m]$ **do**

13: **if** $r_{\mathcal{G}, \alpha, \sigma}(c_i \chi_i) > r_{\mathcal{G}, \alpha, \sigma}(k)$, **then** $k \leftarrow c_i \chi_i$.

14: **end for**

15: **return** $k := (k_1, \dots, k_m)$.

16: **end Procedure**

approximation based on submodular set function maximization. However, our problem is more general in that we want to select multi-sets from $\mathcal{W}$ with budget constraints, and the reward function in general does not have the DR-submodular property for an arbitrary starting distribution. Nevertheless, we have the following lemma to solve our problem.

Lemma 3. For any network $\mathcal{G}$, distribution α and weights σ , $r_{\mathcal{G}, \alpha, \sigma}(\cdot) : \mathbb{Z}_{\geq 0}^m \rightarrow \mathbb{R}$ is monotone and submodular.

Leveraging on the monotone submodular property, we design a Budget Effective Greedy algorithm (Algorithm 1). The core of Algorithm 1 is the procedure of BEG. Let $\delta(i, b, k)$ be the *per-unit marginal gain* $(r_{\mathcal{G}, \alpha, \sigma}(k + b \chi_i) - r_{\mathcal{G}, \alpha, \sigma}(k))/b$ for allocating b more budgets to layer i , which equals to

$$\sum_{u \in \mathcal{V}} \sigma_u \prod_{j \neq i} (1 - P_{j,u}(k_j)) (P_{i,u}(k_i + b) - P_{i,u}(k_i)) / b. \quad (11)$$

The BEG procedure consists of two parts and maintains a queue Q , where any pair $(i, b) \in Q$ represents a tentative plan of allocating b more budgets to layer i . The first part is built around the while loop (lines 6–10), where each iteration greedily selects the (i, b) pair in Q such that the *per-unit marginal gain* $\delta(i, b, k)$ is maximized. If b is within the remaining budget, we allocate b to layer i ; otherwise, we proceed to the next iteration. The second part is a for loop (lines 12–13), where in the i -th round we attempt to allocate all c_i budgets to layer i and replace the current best budget allocation if we have a larger reward.

Theorem 1. Algorithm 1 obtains a $(1 - e^{-\eta}) \approx 0.357$ -approximate solution, where η is the solution of equation $e^{\eta} = 2 - \eta$, to the overlapping MuLaNE problem with an arbitrary starting distribution.

Proof. See Section C.1 for details. $\square$

Line 1 uses $O(m\|c\|_{\infty}n_{\max}^3)$ time to pre-calculate visiting probabilities based on Eq. (9), where $n_{\max} = \max_i |\mathcal{V}_i|$. In the procedure of BEG, the while loop contains B iterations, the size of queue Q is $O(m\|c\|_{\infty})$, and line 7 uses $O(n_{\max})$ to calculate $\delta(i, b, k)$ by proper pre-computation and update, thus the time complexity of Algorithm 1 is $O(B\|c\|_{\infty}mn_{\max} + m\|c\|_{\infty}n_{\max}^3)$.

Remark 1. The idea of combining the greedy algorithm with enumerating solutions on one layer is also adopted in previous works [12, 25], but they only give a $\frac{1}{2}(1 - e^{-1}) \approx 0.316$ -approximation analysis. In this paper, we provide a novel analysis with a better $(1 - e^{-\eta}) \approx 0.357$ -approximation.

Remark 2. Another algorithm proposed by Alon et al. [12], based on partial enumeration techniques (i.e., BEGE), achieves a $(1 - 1/e)$ approximation ratio. This ratio is the best possible in polynomial time unless $P = NP$. However, the algorithm incurs prohibitively high computational cost, with a time complexity of $O(B^4 m^4 \|c\|_{\infty} n_{\max} + m\|c\|_{\infty} n_{\max}^3)$. Similarly, the algorithm in Soma et al. [13] also

Algorithm 2 Myopic greedy (MG) algorithm for the non-overlapping MuLaNE.

```

1: Same input, output and lines 1–2 as in Algorithm 1, except replacing BEG with MG.
2: Procedure MG( $(P_{i,u}(b))_{i \in [m], u \in \mathcal{V}, b \in [c_i]}, \sigma, B, c$ )
3: Let  $k := (k_1, \dots, k_m) \leftarrow \mathbf{0}$ ,  $K \leftarrow B$ .
4: while  $K > 0$  do
5:    $i^* \leftarrow \arg \max_{i \in [m], k_i + 1 \leq c_i} \delta(i, 1, k)$ . {Eq. (11)}
6:    $k_{i^*} \leftarrow k_{i^*} + 1$ ,  $K \leftarrow K - 1$ .
7: end while
8: return  $k = (k_1, \dots, k_m)$ .
9: end Procedure

```

achieves a $(1 - 1/e)$ approximation ratio but with a time complexity of $O(mB^5 n_{\max}^4)$. We also provide the algorithm with a $(1 - 1/e)$ -approximate solution for the overlapping MuLaNE problem with arbitrary starting distributions, as detailed in the analysis in D.

Starting from Stationary Distributions. We also consider the special case where each random explorer W_i starts from the stationary distribution π_i with $\pi_i^\top P = \pi_i^\top$. If a Markov chain P admits a stationary distribution π , the random walk will gradually approach π from arbitrary starting distributions. Furthermore, in the random walk literature, the stationary distribution often gives nice properties for the theoretical analysis. In this case, we have the following stronger DR-submodularity.

Lemma 4. *For any network $\mathcal{G}$, stationary distributions π and node weights σ , function $r_{\mathcal{G}, \pi, \sigma}(\cdot) : \mathbb{Z}_{\geq 0}^m \rightarrow \mathbb{R}$ is monotone and DR-submodular, i.e., $r_{\mathcal{G}}(\pi, y + \chi_j) - r_{\mathcal{G}}(\pi, y) \leq r_{\mathcal{G}}(\pi, x + \chi_j) - r_{\mathcal{G}}(\pi, x)$, and $r_{\mathcal{G}}(\pi, x) \leq r_{\mathcal{G}}(\pi, y)$ for any $x \leq y$, $j \in [m]$.*

Since the reward function is DR-submodular, the procedure of BEG can be replaced by the simple MG procedure in Algorithm 2 with a better approximation ratio. The time complexity is also improved to $O(Bmn_{\max} + m\|c\|_{\infty} n_{\max}^3)$. As shown in Algorithm 2, first, line 5 always select the pair (i^*, b^*) such that $b^* = 1$ because $r_{\mathcal{G}}(\alpha, k + b\chi_i) - r_{\mathcal{G}}(\alpha, k + (b-1)\chi_i) \leq r_{\mathcal{G}}(\alpha, k + \chi_i) - r_{\mathcal{G}}(\alpha, k)$ for arbitrary k , i and b . Thus, Algorithm 1 can always extend the current budget allocation k by 1 and never enters line 10. Therefore, line 10 and lines 12–13 of Algorithm 1 can be safely removed in Algorithm 2. Following a similar argument to the proof of Theorem 1, we can show that the following theorem holds.

Theorem 2. *Algorithm 2 obtains a $(1 - 1/e)$ -approximate solution to the overlapping MulaNE with the stationary starting distributions.*

Proof. See Section C.2 for details. $\square$

5.2. Offline algorithm for non-overlapping MuLaNE

For the non-overlapping MuLaNE setting, we are able to obtain the exact optimal solution rather than merely an approximate one. Specifically, under the stationary starting distribution, a slight modification of the greedy algorithm (Algorithm 2) yields the exact optimum; for arbitrary starting distributions, we further design a dynamic programming algorithm that computes the exact optimal solution (see E).

6. Online learning for MuLaNE

This section addresses the online overlapping and non-overlapping MuLaNE problem, where the structure of the network $\mathcal{G}$, the distribution α , and node weights σ are not known a priori. Instead, the only information revealed to the decision maker includes total budget B , the number of layers m , the number of the target nodes $|\mathcal{V}|$ (or an upper bound of it), and the budget constraint c . For the unknown network structure, our solution is to estimate certain intermediate parameters, instead of estimating the graph structure (namely the transition probability matrix P_i) and the starting distribution α . While one could try to estimate each transition probability matrix P_i directly, this method presents two major issues. First, it is difficult to analyze how the estimated $\hat{P}_i$ affects the performance of online algorithms, since the reward function given by Eq. (7) is highly non-linear. Second, this method requires extensive matrix calculations in each round, which is inefficient for large networks. Subsequently, our solution bypasses the transition matrices P_i and instead directly estimate the visiting probabilities $P_{i,u}(b) \in [0, 1]$. For unknown node weight σ_u , we maintain an optimistic weight before we first visit u . With a little abuse of the notation, we use μ to denote the unknown visiting probability parameters and $r_{\mu, \sigma}(k)$ to denote the reward $r_{\mathcal{G}, \alpha, \sigma}(k)$.

Specifically, we maintain a set of base arms $\mathcal{A} = \{(i, u, b) | i \in [m], u \in \mathcal{V}, b \in [c_i]\}$, where the total number $|\mathcal{A}| = \sum_{i \in [m]} c_i |\mathcal{V}|$. For each base arm $(i, u, b) \in \mathcal{A}$, we denote $\mu_{i,u,b}$ as the true value of each base arm, i.e., $\mu_{i,u,b} = P_{i,u}(b)$. For the unknown fixed node weights, we maintain the optimistic weight $\bar{\sigma}_u = 1$ if $u \in \mathcal{V}$ has not been visited. After u is first visited and its true value σ_u is revealed, we replace $\bar{\sigma}_u$ with σ_u . Beyond the fixed-node-weight scenario, we also consider the more complex case of random node weights, which will be detailed in subsequent sections.

To provide a high-level overview of our online learning framework, we present the complete workflow in Fig. 4. The process is an iterative feedback loop designed to operate on an initially unknown network structure. In each round, the agent begins with its

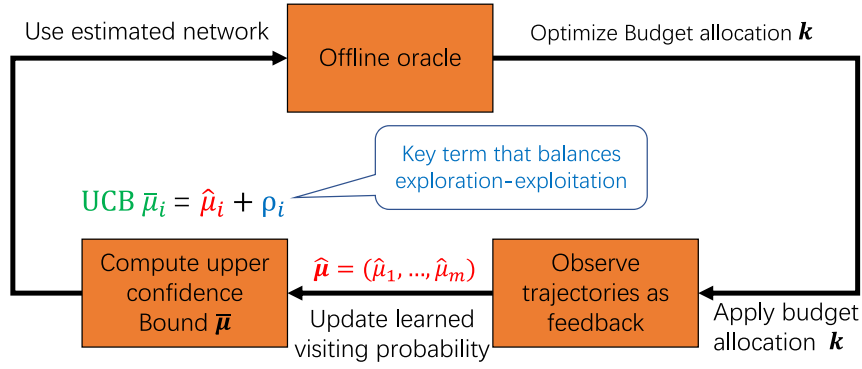


Fig. 4. Complete iterative process of MuLaNE.

current estimates of the network parameters (e.g., the visiting probabilities). It first calls an offline oracle to determine an optimized budget allocation k , based on these estimates. The agent then executes the exploration by performing random walks according to this budget and observes the resulting trajectories as feedback. This new information is used to update and refine the learned parameters for the next round. This cycle of budget optimization, exploration, observation, and model refinement continues for T rounds, with the overarching goal of maximizing the total cumulative reward by progressively learning a more accurate representation of the network's dynamics.

6.1. Online algorithm for overlapping MuLaNE

We present our algorithm in Algorithm 3, which is an adaptation of the CUCB algorithm for the general Combinatorial Multi-arm Bandit (CMAB) framework [7,26] to our setting. Notice that the set of base arms $\mathcal{A}$ is defined using $\mathcal{V}$, the set of node IDs, which is not known before the learning process begins. This is handled by operating on a known size of the node set $|\mathcal{V}|$. At the start, the algorithm creates $|\mathcal{V}|$ placeholders for the node IDs. As new, previously unseen nodes are visited during exploration, they are dynamically assigned to these placeholders. Therefore, the use of $\mathcal{V}$ in the definition of $\mathcal{A}$ is for notational convenience and refers to this dynamic placeholder set.

In Algorithm 3, we maintain an unbiased estimation of the visiting probability $P_{i,u}(b)$, denoted as $\hat{\mu}_{i,u,b}$. Let $T_{i,u,b}$ record the number of times arm (i, u, b) is played so far and $\bar{\sigma}_v$ denotes the optimistic importance weight. In each round $t \geq 1$, we compute the radius $\rho_{i,u,b}$ in line 5, which controls the level of exploration. The confidence radius is larger when the arm (i, u, b) is not explored often (i.e. $T_{i,u,b}$ is small), and thus motivates more exploration. In line 6, we compute the upper confidence bound (UCB) value $\tilde{\mu}_{i,u,b}$. Due to the randomness of the exploration process, the upper confidence bound (UCB) value $\tilde{\mu}_{i,u,b}$ could be decreasing w.r.t b , but our offline oracle BEG can only accept non-decreasing UCB values (otherwise the $1 - e^{-\eta}$ approximation is not guaranteed). Therefore, in line 8, we increase the UCB value $\tilde{\mu}_{i,u,b}$ and set it to be $\max_{j \in [b]} \tilde{\mu}_{i,u,j}$. This is the adaption of the CUCB algorithm to fix our case, and thus we name our algorithm as CUCB-MAX. Consider two base arms $a_1 = (i, u, b_1)$ and $a_2 = (i, u, b_2)$, $b_1 < b_2$, arm a_2 has both larger true value and larger confidence radius (T_{i,u,b_2} is smaller due to line 13), so it is safe to increase a_2 's UCB value if it is smaller than a_1 's. After we apply the $1 - e^{-\eta}$ approximate solution k given by the oracle (i.e., Algorithm 1), we get m trajectories as the feedback. In line 11, we can update the unknown weights of visited nodes. In line 13, for base arms (i, u, b) with $b \leq k_i$, we update corresponding statistics by the Bernoulli random variable $Y_{i,u,b} \in \{0, 1\}$ indicating whether node u is visited by W_i in first b steps.

Regret Bound Analysis. To state the regret bound, we first give some definitions followed by our main result. We define the reward gap Δ_k as $\max(0, \xi r_{\mu,\sigma}(k^*) - r_{\mu,\sigma}(k))$ for all feasible action k satisfying $\sum_{i=1}^m k_i = B$, $0 \leq k_i \leq c_i$, where k^* is the optimal solution for parameters μ, σ and $\xi = 1 - e^{-\eta}$ is the approximation ratio of the $(\xi, 1)$ -approximate offline oracle. For each base arm (i, u, b) , we define $\Delta_{\min}^{i,u,b} = \min_{\Delta_k > 0, k_i = b} \Delta_k$ and $\Delta_{\max}^{i,u,b} = \max_{\Delta_k > 0, k_i = b} \Delta_k$. As a convention, if there is no action k with $k_i = b$ such that $\Delta_k > 0$, we define $\Delta_{\min}^{i,u,b} = \infty$ and $\Delta_{\max}^{i,u,b} = 0$. Let $\Delta_{\min} = \min_{(i,u,b) \in \mathcal{A}} \Delta_{\min}^{i,u,b}$ and $\Delta_{\max} = \max_{(i,u,b) \in \mathcal{A}} \Delta_{\max}^{i,u,b}$.

Theorem 3. Algorithm 3 has the following distribution-dependent $(1 - e^{-\eta}, 1)$ approximation regret, $O\left(\sum_{(i,u,b) \in \mathcal{A}} \frac{m|\mathcal{V}| \log T}{\Delta_{\min}^{i,u,b}}\right)$.

Proof. See Section F.2 for details. $\square$

Remark 1. Looking at the above distribution-dependent bound, we have the $O(\log T)$ approximation regret, which is asymptotically tight. Coefficient $m|\mathcal{V}|$ in the leading term corresponds to the number of edges in the complete bipartite coverage graph. Notice that we cannot use the true edge number $\sum_{i \in [m]} |\mathcal{V}_i|$, because the learning algorithm does not know which nodes are contained in each layer, and has to explore all visiting possibilities given by the default complete bipartite graph. For the non-overlapping case, we will remove the coefficient $m|\mathcal{V}|$ to achieve a tighter regret bound.

Remark 2. The $(1 - e^{-\eta}, 1)$ -approximate regret arises from the offline oracle BEG used in line 9. This can be improved to $(1 - 1/e, 1)$ regret by employing BEGE, or even to the exact regret if the oracle computes the exact optimal budget allocation. The usage of BEG is a trade-off we make between computational efficiency and learning efficiency, and empirically, it performs well as we shall see in Section 7.

Algorithm 3 CUCB-MAX algorithm for the MuLaNE.

Input: Budget B , number of layers m , number of nodes $|\mathcal{V}|$, constraints c , offline oracle BEG.

- 1: For each arm $(i, u, b) \in \mathcal{A}$, $T_{i,u,b} \leftarrow 0$, $\hat{\mu}_{i,u,b} \leftarrow 0$.
- 2: For each node $v \in \mathcal{V}$, $\bar{\sigma}_v \leftarrow 1$.
- 3: **for** $t = 1, 2, 3, \dots, T$ **do**
- 4: **for** $(i, u, b) \in \mathcal{A}$ **do**
- 5: $\rho_{i,u,b} \leftarrow \sqrt{3 \ln t / (2T_{i,u,b})}$.
- 6: $\tilde{\mu}_{i,u,b} \leftarrow \min\{\hat{\mu}_{i,u,b} + \rho_{i,u,b}, 1\}$.
- 7: **end for**
- 8: For $(i, u, b) \in \mathcal{A}$, $\tilde{\mu}_{i,u,b} \leftarrow \max_{j \in [b]} \tilde{\mu}_{i,u,j}$.
- 9: $\mathbf{k} \leftarrow \text{BEG}((\tilde{\mu}_{i,u,b})_{(i,u,b) \in \mathcal{A}}, (\bar{\sigma}_v)_{v \in \mathcal{V}}, B, c)$.
- 10: Apply budget allocation $\mathbf{k}$, which gives trajectories $\mathbf{X} := (X_{i,1}, \dots, X_{i,k_i})_{i \in [m]}$ as feedbacks.
- 11: For any visited node $v \in \bigcup_{i \in [m]} \{X_{i,1}, \dots, X_{i,k_i}\}$, receive its node weight σ_v and set $\bar{\sigma}_v \leftarrow \sigma_v$.
- 12: For any $(i, u, b) \in \tau := \{(i, u, b) \in \mathcal{A} \mid b \leq k_i\}$,
 $Y_{i,u,b} \leftarrow 1$ if $u \in \{X_{i,1}, \dots, X_{i,b}\}$ and 0 otherwise.
- 13: For $(i, u, b) \in \tau$, update $T_{i,u,b}$ and $\hat{\mu}_{i,u,b}$:
 $T_{i,u,b} \leftarrow T_{i,u,b} + 1$, $\hat{\mu}_{i,u,b} \leftarrow \hat{\mu}_{i,u,b} + (Y_{i,u,b} - \hat{\mu}_{i,u,b}) / T_{i,u,b}$.
- 14: **end for**

Remark 3. The full proof of the above theorem is included in the F, where we rely on the following properties of $r_{\mu,\sigma}(\mathbf{k})$ to bound the regret. We further consider random node weights with unknown mean vector σ (Please refer to F.3).

Property 1. (Monotonicity). The reward $r_{\mu,\sigma}(\mathbf{k})$ is monotonically increasing, i.e., for any budget allocation $\mathbf{k}$, any two vectors $\mu = (\mu_{i,u,b})_{(i,u,b) \in \mathcal{A}}$, $\mu' = (\mu'_{i,u,b})_{(i,u,b) \in \mathcal{A}}$ and any node weights σ, σ' , we have $r_{\mu,\sigma}(\mathbf{k}) \leq r_{\mu',\sigma'}(\mathbf{k})$, if $\mu_{i,u,b} \leq \mu'_{i,u,b}$, $\sigma_v \leq \sigma'_v$, $\forall (i, u, b) \in \mathcal{A}, v \in \mathcal{V}$.

Property 2. (1-Norm Bounded Smoothness). The reward function $r_{\mu,\sigma}(\mathbf{k})$ satisfies the 1-norm bounded smoothness condition, i.e., for any budget allocation $\mathbf{k}$, any two vectors $\mu = (\mu_{i,u,b})_{(i,u,b) \in \mathcal{A}}$, $\mu' = (\mu'_{i,u,b})_{(i,u,b) \in \mathcal{A}}$ and any node weights σ, σ' , we have $|r_{\mu,\sigma}(\mathbf{k}) - r_{\mu',\sigma'}(\mathbf{k})| \leq \sum_{i \in [m], u \in \mathcal{V}, b=k_i} (\sigma_u |\mu_{i,u,b} - \mu'_{i,u,b}| + |\sigma_u - \sigma'_u| \mu'_{i,u,b})$.

Analysis Technique. We emphasize that our algorithm and analysis differ from the original CUCB algorithm [7] as follows. First, we have the additional regret caused by the over-estimated weights for unvisited nodes, i.e., $|\sigma_u - \sigma'_u| \mu'_{i,u,b}$ term in Property 2. We carefully bound this term on the basis of the observation that $\mu'_{i,u,b}$ is small and rapidly decreases before u is first visited. Next, we have to take the maximum (line 8) to guarantee that the UCB value is monotone w.r.t b since our BEG oracle can only output the $(1 - e^{-\eta}, 1)$ approximation with monotone input. Due to the above operation, the value of the (i, u, b) UCB depends on the feedback from all arms (i, u, j) for $j \leq b$ (set τ in line 12). So we should update all these arms (line 13) to guarantee that the estimates for all these arms are accurate enough. Finally, directly following the standard CMAB result would have a greater regret upper bound, because arms in τ are defined as triggered arms, but only arms in $\tau' = \{(i, u, b) \in \mathcal{A} \mid k_i = b\}$ affect the rewards. So we conceptually view arms in τ' as triggered arms and use a tighter 1-Norm Bounded Smoothness condition as given above to derive a tighter regret bound. This improves the coefficient of the leading $\ln T$ term in the distribution-dependent regret bound by a factor of $|\tau|/|\tau'| = O(B/m)$, and the $1/\Delta_{\min}^{i,u,b}$ term is smaller since the original definition would have $\Delta_{\min}^{i,u,b} = \min_{\Delta_k > 0, b \leq k_i} \Delta_k$.

6.2. Explorer-relaxed online algorithm for overlapping MuLaNE

As shown in Algorithm 4, we now introduce an explorer-relaxed online algorithm called ECUCB-MAX. By “explorer-relaxed”, we mean an algorithm whose regret scales sub-linearly in the number of explorers m , rather than linearly as in previous work, making it substantially more scalable for large network layers. The overall structure of Algorithm 4 is similar to Algorithm 3. However, in addition to maintaining the empirical unbiased estimate $\hat{\mu}_{i,u,b}$ for the visiting probability $P_{i,u}(b)$, Algorithm 4 also tracks the empirical variance estimate $\hat{V}_{i,u,b}$ for each tuple (i, u, b) . In contrast to Algorithm 3, which constructs the confidence interval $\rho_{i,u,b} = \sqrt{\frac{3 \log t}{2T_{i,u,b}}}$ using the Chernoff-Hoeffding inequality, ECUCB-MAX employs the empirical variance $\hat{V}_{i,u,b}$ to form a more refined confidence interval based on the Bernstein inequality:

$$\rho_{i,u,b} = \sqrt{\frac{6\hat{V}_{i,u,b} \log t}{T_{i,u,b}}} + \frac{9 \log t}{T_{i,u,b}}. \quad (12)$$

For the first term in Eq. (12), $\hat{V}_{i,u,b}$ is approximately equal to the true variance $V_{i,u,b}$. With $V_{i,u,b} \leq (1 - \mu_{i,u,b})\mu_{i,u,b}$ ⁴, the estimation of $\mu_{i,u,b}$ becomes more precise when $\mu_{i,u,b}$ approaches 0 or 1, effectively mitigating the influence of $|\mathcal{V}|$, which will be discussed in detail

⁴ For a bounded random variable $X \in [0, 1]$ with mean μ , variance $V = \mathbb{E}[X^2] - \mathbb{E}[X]^2 \leq \mathbb{E}[X] - (\mathbb{E}[X])^2 \leq (1 - \mu)\mu$, with equality for a Bernoulli random variable.

Algorithm 4 ECUCB-MAX algorithm for the MuLaNE.

Input: Budget B , number of layers m , number of nodes $|\mathcal{V}|$, constraints c , offline oracle BEG.

- 1: For each arm $(i, u, b) \in \mathcal{A}$, $T_{i,u,b} \leftarrow 0$, $\hat{\mu}_{i,u,b} \leftarrow 0$.
- 2: For each node $v \in \mathcal{V}$, $\bar{\sigma}_v \leftarrow 1$.
- 3: **for** $t = 1, 2, 3, \dots, T$ **do**
- 4: **for** $(i, u, b) \in \mathcal{A}$ **do**
- 5: $\rho_{i,u,b} \leftarrow \sqrt{\frac{6\hat{V}_{i,u,b} \log t}{T_{i,u,b}}} + \frac{9 \log t}{T_{i,u,b}}$.
- 6: $\tilde{\mu}_{i,u,b} \leftarrow \min\{\hat{\mu}_{i,u,b} + \rho_{i,u,b}, 1\}$.
- 7: **end for**
- 8: For $(i, u, b) \in \mathcal{A}$, $\bar{\mu}_{i,u,b} \leftarrow \max_{j \in [b]} \tilde{\mu}_{i,u,j}$.
- 9: $\mathbf{k} \leftarrow \text{BEG}((\bar{\mu}_{i,u,b})_{(i,u,b) \in \mathcal{A}}, (\bar{\sigma}_v)_{v \in \mathcal{V}}, B, c)$.
- 10: Apply budget allocation $\mathbf{k}$, which gives trajectories $\mathbf{X} := (X_{i,1}, \dots, X_{i,k_i})_{i \in [m]}$ as feedbacks.
- 11: For any visited node $v \in \bigcup_{i \in [m]} \{X_{i,1}, \dots, X_{i,k_i}\}$, receive its node weight σ_v and set $\bar{\sigma}_v \leftarrow \sigma_v$.
- 12: For any $(i, u, b) \in \tau := \{(i, u, b) \in \mathcal{A} \mid b \leq k_i\}$,
 $Y_{i,u,b} \leftarrow 1$ if $u \in \{X_{i,1}, \dots, X_{i,b}\}$ and 0 otherwise.
- 13: For every $(i, u, b) \in \tau_t$, update $T_{i,u,b}$, $\hat{V}_{i,u,b}$ and $\hat{\mu}_{i,u,b}$: $T_{i,u,b} \leftarrow T_{i,u,b} + 1$, $\hat{\mu}_{i,u,b} \leftarrow \hat{\mu}_{i,u,b} + (X_{i,u,b} - \hat{\mu}_{i,u,b})/T_{i,u,b}$, $\hat{V}_{i,u,b} \leftarrow \frac{T_{i,u,b}-1}{T_{i,u,b}} \left(\hat{V}_{i,u,b} + \frac{1}{T_{i,u,b}} (\hat{\mu}_{i,u,b} - X_{i,u,b})^2 \right)$.
- 14: **end for**

later. The second term of Eq. (12) compensates for using the empirical variance $\hat{V}_{i,u,b}$ instead of the unknown true variance $V_{i,u,b}$, which will be shown in G.2.1.

Regret Bound Analysis. Following the definition in Section 6.1, we have the following theorem.

Theorem 4. Algorithm 4 has the following distribution-dependent $(1 - e^{-\eta}, 1)$ approximation regret, $O\left(\sum_{(i,u,b) \in \mathcal{A}} \left(\frac{|\mathcal{V}| \log(m|\mathcal{V}) \log T}{\Delta_{i,u,b}^{\min}} + \log\left(\frac{m|\mathcal{V}|}{\Delta_{i,u,b}^{\min}}\right) \log T\right)\right)$.

Proof. See Section G.2 for details. $\square$

Remark 1. Taking into account the upper bound of regret above, the dependence on the number of explorers m is $\tilde{O}(\log m)$. In comparison to Theorems 3 and 4 provides a strictly better regret bound, with an improvement factor of $O\left(\frac{m}{\log m|\mathcal{V}|}\right)$.

Remark 2. In comparison to Wang and Chen [17], our results match theirs up to a factor of $O(\log m|\mathcal{V}|)$. Regarding the lower bound, based on the findings of Merlis and Mannor [27], our regret bound remains tight within the same factor. Compared to Liu et al. [28], we have removed a factor $O(\log m|\mathcal{V}|)$ from the second main term in the regret upper bound by effectively reducing the influence of observation randomness.

We then introduce a new property discovered for the MuLaNE problem based on Algorithm 4.

Property 3. (Precision-Balanced Bounded Smoothness). The reward function $r_{\mu,\sigma}(\mathbf{k})$ satisfies the precision-balanced bounded smoothness condition, i.e., for any budget allocation $\mathbf{k}$, any two vectors $\mu = (\mu_{i,u,b})_{(i,u,b) \in \mathcal{A}}$, $\mu' = \mu + \zeta_\mu + \eta_\mu = (\mu'_{i,u,b})_{(i,u,b) \in \mathcal{A}}$ and any node weights $\sigma, \sigma' = \sigma + \zeta_\sigma + \eta_\sigma$, where $\zeta, \eta \in [-1, 1]^m$ represent the parameter change, we have the inequality $|r_{\mu,\sigma}(\mathbf{k}) - r_{\mu',\sigma'}(\mathbf{k})| \leq$

$$\sqrt{1.25|\mathcal{V}|} \sqrt{\sum_{u \in \mathcal{V}} \left(\sum_{i \in [m], b=k_i} \left(\frac{\zeta_{\mu,i,u}^2}{(1-\mu_{i,u,b})\mu_{i,u,b}} \right) + \frac{\zeta_{\sigma,u}^2}{(1-\sigma_u)\sigma_u} \right) + \sum_{u \in \mathcal{V}} \left(\sum_{i \in [m], b=k_i} |\eta_{\mu,i,u,b}| + |\eta_{\sigma,u}| \right)}.$$

Remark 3. Although the impact of $|\mathcal{V}|$ can increase significantly with abrupt reward changes when $\mu_{i,u,b}$ approaches 0 or 1, introducing the term $1/(1 - \mu_{i,u,b})\mu_{i,u,b}$ amplifies the square root term in these cases, thereby mitigating the negative effect. Moreover, as $\mu_{i,u,b}$ nears 0 or 1, the variance $V_{i,u,b}$ diminishes, allowing for “a more precise estimation” of $\mu_{i,u,b}$. This high estimation precision leads to a smaller estimation gap, introducing a variance-related term that offsets the challenging $(1 - \mu_{i,u,b})\mu_{i,u,b}$ term in the denominator. This reveals an inherent balance: regions where the reward function is less smooth are precisely where its underlying parameters can be estimated with higher precision. Since this property captures the *balance between estimation precision and the smoothness bound*, we refer to it as “Precision-Balanced Bounded Smoothness”.

Analysis Technique. In addition to the analysis technique in Section 6.1, the proof for Theorem 4 is more challenging than that for Theorem 3. Drawing inspiration from Liu et al. [28], our approach filters the total regret through several events and bounds the regret for each filtered event separately. As shown in G.2, the event contributing most to the regret is $E_t = \{\Delta_{k_t} \leq e_t(k_t)\}$, where

the error term $e_t(k_t) = O(|\mathcal{V}| \sqrt{\sum_{u \in \mathcal{V}} \sum_{i \in [m], b=k_i} \left(\frac{\log t}{T_{i,u,b}} \right) + \sum_{u \in \mathcal{V}} \sum_{i \in [m], b=k_i} \left(\frac{\log t}{T_{i,u,b}} \right)})$. To bound the leading regret filtered by E_t , we use the reverse amortization trick [17,29], and adaptively allocate each arm’s regret contribution, carefully selecting thresholds for the error term $e_t(k_t)$ specific to the MuLaNE problem rather than relying on those in Wang and Chen [17], Liu et al. [28]. Additionally, our analysis holds for both the primary model of “fixed but unknown” node weights and the more general stochastic model where weights are “random variables” with an unknown mean (as detailed in G). In both settings, we avoid the triggering probability group technique from Wang and Chen [17], Liu et al. [28] and instead minimize observation uncertainty.

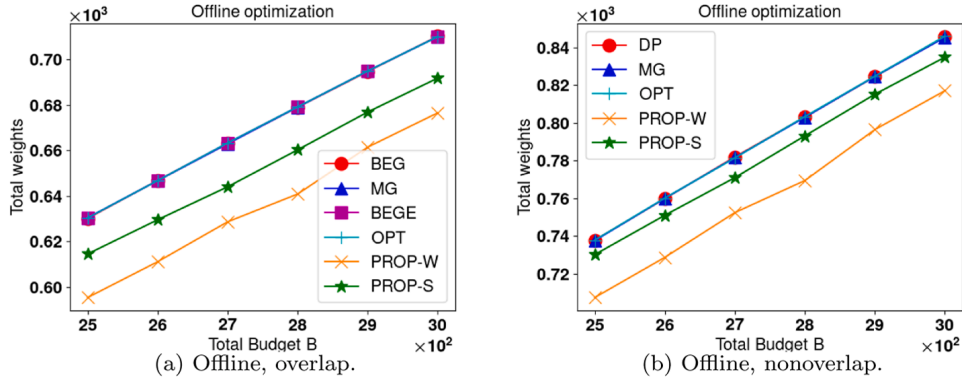


Fig. 5. Total weights of unique nodes visited for offline algorithms.

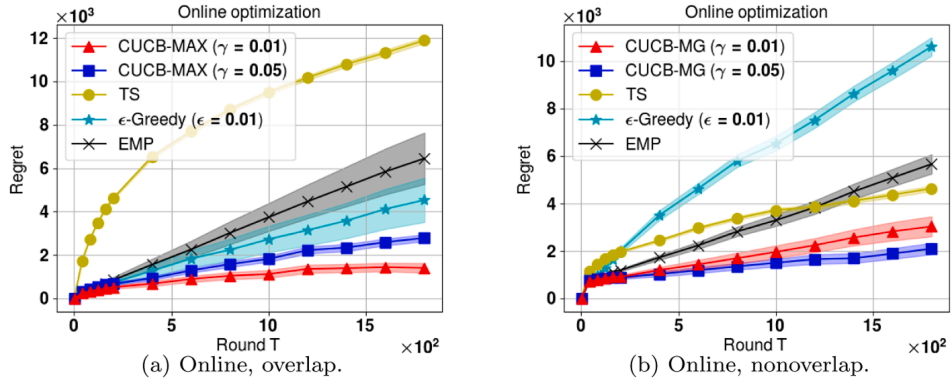
Fig. 6. Cumulative regret for online algorithms when $B = 3000$.

Table 2
Statistics for FF-TW-YT network.

Layer	FriendFeed	Twitter	YouTube
# of vertices	5540	5702	663
# of edges	31,921	42,327	614

6.3. Online algorithm for non-overlapping case

For the non-overlapping case, we set *layer-wise marginal gains* as our base arms. Concretely, we maintain a set of base arms $\mathcal{A} = \{(i, b) | i \in [m], b \in [c_i]\}$. For each base arm $(i, b) \in \mathcal{A}$, let $\mu_{i,b} = \sum_{u \in \mathcal{V}} \sigma_u(P_{i,u}(b) - P_{i,u}(b-1))$ be the true marginal gain of assigning budget b in layer i . We apply the standard CUCB algorithm to this setting and call the resulting algorithm CUCB-MG (Algorithm 8 in H). Note that in the non-overlapping setting, we can solve the offline problem exactly so that we can use the exact offline oracle to solve the online problem and achieve an exact regret bound. This is the major advantage over the overlapping setting where we can only achieve an approximate bound. Define $\Delta_k = r_{\mu,\sigma}(k^*) - r_{\mu,\sigma}(k)$ for all feasible action k , and $\Delta_{\min}^{i,b} = \min_{\Delta_k > 0, k_i \geq b} \Delta_k$, CUCB-MG has a $O(\sum_{(i,b) \in \mathcal{A}} 48B \ln T / \Delta_{\min}^{i,b})$ regret upper bound.

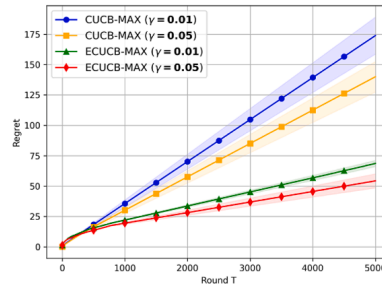
7. Experiments

Experimental Settings. We firstly conduct experiments on a real-world multi-layered network FF-TW-YT dataset, which contains $m = 3$ layers representing users' social connections in FriendFeed (FF), Twitter (TW) and YouTube (YT) [30]. This multi-layered network is built upon the FF network, and the TW and YT layers are generated by exploiting the feature of FriendFeed as a social media aggregator. In total, FF-TW-YT has 6,407 distinct vertices representing users and 74,836 directed edges representing connections ("who follows whom") among users. The statistics for each layer is summarized in Table 2. We transform the FF-TW-YT network $\mathcal{G}(\mathcal{V}, \mathcal{E})$ into a symmetric directed network (by adding a new edge (v, u) if $(u, v) \in \mathcal{E}$ but $(v, u) \notin \mathcal{E}$) because a user can be visited via her followers or followees. Each edge weight is set to be 1 and the node weights are set to be $\sigma_u \in \{0, 0.5, 1\}$ uniformly at random. Each random walker always starts from the smallest node-id in each layer and we set constraints c_i equal to the total budget B . Note that in order to test the non-overlapping case, we use the same network but relabel node-ids so that they do not overlap between different layers.

Table 3Running time (seconds) for offline algorithms with different budgets B .

(a) Offline, overlapping case.			
Algorithm	B = 2.6k	B = 2.8k	B = 3.0k
BEG	0.274	0.316	0.363
BEGE	34.37	45.51	59.20
OPT	91.16	98.42	105.63

(b) Offline, non-overlapping case.			
Algorithm	B = 2.6k	B = 2.8k	B = 3.0k
DP	0.038	0.044	0.050
MG	0.008	0.009	0.010
OPT	70.10	75.78	81.11

**Fig. 7.** Comparison of online algorithms with same offline algorithm in the complex overlapping case.

Similarly, we utilize the RocketFuel dataset [31], which represents a real-world computer network. Our analysis focuses on a randomly selected network from this dataset, comprising 300 nodes and 549 edges. We calculate and normalize the centrality of the interconnectivity for each node to represent the edge weight, a metric that indicates the functional importance of the node within the network [32]. A node's capacity is positively correlated with its functional importance; the greater the capacity, the higher the likelihood that the node remains operational. Identifying such nodes is crucial for a random walker aiming to maximize the discovery of important nodes, with m also set to 3. To handle the randomness, we repeat 10,000 times and present the averaged total weights of the unique nodes visited for offline optimization. We calculate the cumulative regret by comparing with the *optimal* solution, which is stronger than comparing with the $(1 - e^{-\eta}, 1)$ -approximate solution as defined in Eq. (4).

Benchmarks. For the offline setting, we present the results for Algorithm 1 (denoted as BEG), Algorithm 2 (denoted as MG), BEG with partial enumeration (denoted as BEGE) and the dynamic programming algorithm in the supplementary material (denoted as DP). We provide two baselines PROP-S and PROP-W, which allocate the budget proportional to the layer size and proportional to the total weights if we allocate $B/3$ budgets to that layer, respectively. The optimal solution (denoted as OPT) is also provided by enumerating all possible budget allocations. For online settings, we consider CUCB-MAX (Algorithm 3), ECUCB-MAX (Algorithm 4) for the overlapping case, and CUCB-MG for the non-overlapping case (Algorithm 8). We shrink the confidence interval by γ , i.e., $\rho_{i,u,b} \leftarrow \gamma \rho_{i,u,b}$ to speed up the learning, though our theoretical regret bound requires $\gamma = 1$. For baselines, we consider the EMP algorithm which always allocates according to the empirical mean, and the ϵ -Greedy algorithm which allocates budgets according to empirical mean with probability $1 - \epsilon$ and allocates all B budgets to the i -th layer with probability ϵ/m . We also compare with the Thompson sampling (TS) method [29], which uses Beta distribution $\text{Beta}(\alpha, \beta)$ (where $\alpha = \beta = 1$ initially) as prior distribution for each base arm.

Experimental Results. We first present the results of the social network dataset for the MuLaNE problem under the offline setting in Fig. 5. For the offline overlapping case, both BEG and MG outperform two baselines PROP-W and PROP-S in receiving total weights. Although not guaranteed by the theory, BEG is empirically close to BEGE and the optimal solution (OPT). For computational efficiency, in Table 3a, BEG is at least two orders of magnitude (e.g., 163 times when $B = 3.0k$) faster than BEGE and OPT. Combining the result that the reward of BEG is empirically close to BEGE and the optimal solution, this shows that BEG is empirically better than BEGE and OPT. For the offline non-overlapping case, the results presented in Table 3b are similar; however, a key distinction is that we provide a theoretical guarantee for the optimality of DP.

For the online setting, as shown in Fig. 6, all of our proposed algorithms consistently outperform the baseline methods. For the online overlapping case, the curve of CUCB-MAX grows sub-linearly w.r.t round T under the social network dataset, because our offline oracle achieves near-optimal solutions. Moreover, the cumulative regret decreases and consistently outperforms baseline algorithms when $\gamma = 0.1, 0.05$. The resulting sub-linear cumulative regret curve demonstrates empirically that the CUCB-MAX algorithm is effectively learning to make near-optimal budget allocation decisions while optimizing the objective. Moreover, as shown in Fig. 7 on the computer network dataset with the offline algorithm, ECUCB-MAX exhibits lower regret compared to CUCB-MAX with the same

number of random walkers, and shows less fluctuation due to the consideration of variance. In the non-overlapping case, the CUCB-MG algorithm outperforms both baseline approaches. Additional experimental results, including variations in budgets, stationary starting distributions, and the computational efficiency of different online learning algorithms, can be found in [I](#).

8. Conclusion

This paper explores the multi-layered networks through random walks (MuLaNE), where the objective is to maximize the total weight of distinct nodes visited across multiple layers. We systematically study the MuLaNE problem in both offline optimization and online learning settings, considering both overlapping and non-overlapping networks. For the offline setting, we propose four algorithms based on the network's structure (overlapping or non-overlapping) and the nature of the starting distributions (arbitrary or stationary). Each algorithm is accompanied by provable guarantees on approximation factors and running time. In the online setting, where the network structure and node weights are unknown, we develop two distinct online learning algorithms for the more complex overlapping MuLaNE, while also addressing the non-overlapping case. All of these algorithms achieve a logarithmic order regret bound on the total rounds. Finally, we validate the empirical performance of our algorithms through experiments on real-world social and computer network datasets.

Several promising directions remain for future work. One interesting extension would be to allow joint optimization of the starting distributions and budget allocations. Additionally, exploring adaptive versions of MuLaNE, where feedback from earlier exploration informs future strategies, presents a compelling avenue for further research.

CRedit authorship contribution statement

Xiangxiang Dai: Writing – original draft, Visualization, Validation, Methodology, Investigation, Formal analysis, Conceptualization; **Xutong Liu:** Writing – review & editing, Writing – original draft, Validation, Formal analysis, Conceptualization; **Jinhang Zuo:** Investigation, Formal analysis, Conceptualization; **Xiaowei Chen:** Methodology, Conceptualization; **Wei Chen:** Supervision, Methodology, Formal analysis; **John C S Lui:** Writing – review & editing, Supervision, Funding acquisition.

Data availability

This paper primarily makes theoretical contributions. The validation dataset used are open-source datasets, which have been specifically referenced in the paper.

Declaration of competing interests

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgements

The work of Xiangxiang Dai was supported in part by the [National Natural Science Foundation of China \(625B2163\)](#). The work of John C.S. Lui was supported in part by the RGC GRF-14202923.

Appendix A. Appendix overview

The appendix is organized as follows. We provide proofs and examples for properties of the visiting probability in [B](#). Proofs of offline optimization for overlapping MuLaNE are provided in [C](#). We provide the detailed budget-effective greedy algorithm with partial enumeration (BEGE) and its analysis in [D](#). Proofs of offline optimization for non-overlapping MuLaNE are provided in [E](#). We state the detailed analysis of online learning and explorer-relaxed online learning for overlapping MuLaNE in [F](#) and [G](#), respectively. We state the analysis of online learning for non-overlapping MuLaNE in [H](#). Supplemental experiments are provided in [I](#).

Appendix B. Proofs and examples for properties of visiting probability

B.1. Starting from arbitrary distributions

Proof. For analysis, we use the following equivalent formulation for $k_i \geq 2$, (trivially $P_{i,u}(0) = 0$ and $P_{i,u}(1) = \alpha_{i,u}$),

$$P_{i,u}(k_i) = ((\alpha_i)_{-u}^\top (\Gamma_i)_{uk_i} + \alpha_{i,u}), \quad (\text{B.1})$$

where $(\Gamma_i)_{uk_i} = (I + T_i(u) + \dots + T_i(u)^{k_i-2})(p_i)_u$, $(p_i)_u = (P_i[\cdot, u])_{-u}$, $T_i(u) = (P_i)_{-u, -u}$, (we use Eq. (2.67) in Kijima [\[33\]](#) to derive Γ_{uk_i}). Note that $p_{-u} \in \mathbb{R}^{n-1}$ is the vector obtained by deleting the u -th element from $p \in \mathbb{R}^n$, and $P_{-u, -v} \in \mathbb{R}^{(n-1) \times (n-1)}$ is the matrix obtained by deleting the u -th row and the v -th column from $P \in \mathbb{R}^{n \times n}$.

We further derive the marginal gain of $P_{i,u}(\cdot)$ at step $k_i \geq 2$ (trivially $g_{i,u}(1) = \alpha_{i,u}$) as,

$$g_{i,u}(k_i) = (\alpha_i)_{-u}^\top T_i(u)^{k_i-2} (p_i)_u. \quad (\text{B.2})$$

Now, we can show $g_{i,u}(k_i)$ is non-negative because any element of $(\alpha_i)_{-u}^\top$, $T_i(u)$ and $(p_i)_u$ are non-negative, which means $P_{i,u}(k_i)$ is non-decreasing with respect to step k_i . $\square$

An example showing $g_{i,u}(k_i)$ is not monotone: Consider a path P_3 with three nodes as the i -th layer $\mathcal{G}_i(\mathcal{V}_i, \mathcal{E}_i)$, where $\mathcal{V}_i = \{u, v, w\}$ and $\mathcal{E}_i = \{(u, v), (v, u), (v, w), (w, v)\}$. If the W_i always starts from the left-most node u and chooses the right-most node w as our target node. Then, $g_{i,w}(1) = g_{i,w}(2) = 0$ but $g_{i,w}(3) > 0$ since at least three steps are needed to visit the node w , which shows that $g_{i,w}(k_i)$ is not always non-increasing.

B.2. Starting from the stationary distribution

Proof. Consider any layer L_i with transition probability matrix P_i , if we start from the stationary distribution π_i , the probability that node $u \in \mathcal{V}_i$ is ever visited in the first k_i steps is $P_{i,u}(k_i) = (\pi_i)_{-u}^\top (\Gamma_i)_{u k_i} + \pi_{i,u}$, where $(\Gamma_i)_{u k_i} = (I + T_i(u) + \dots + T_i(u)^{k_i-2})(p_i)_u$, $T_i(u) = (P_i)_{-u, -u}$, $(p_i)_u = (P_i[\cdot, u])_{-u}$. Then the marginal gain for node u is $g_{i,u}(k_i) = P_{i,u}(k_i) - P_{i,u}(k_i - 1) = (\pi_i)_{-u}^\top T_i(u)^{k_i-2} (p_i)_u$ for $k_i \geq 2$ and $g_{i,u}(k_i) = \pi_{i,u}$ when $k_i = 1$.

Define the margin of the marginal gain as $\Delta_{i,u}(k_i) = g_{i,u}(k_i + 1) - g_{i,u}(k_i)$. When $k = 1$, $\Delta(u, 1) = g_{i,u}(2) - g_{i,u}(1) = (\pi_i)_{-u}^\top \cdot (p_i)_u - \pi_{i,u} = \pi_{i,u} - \pi_{i,u} P[u, u] - \pi_{i,u} = -\pi_{i,u} P[u, u] \leq 0$. When $k_i \geq 2$, $\Delta_{i,u}(k_i) = g_{i,u}(k_i + 1) - g_{i,u}(k_i) = (\pi_i)_{-u}^\top (T_i(u)^{k_i-1} - T_i(u)^{k_i-2})(p_i)_u$. Because $\pi_i^\top P_i = \pi_i^\top$, we have $(\pi_i)_{-u}^\top - (\pi_i)_{-u}^\top P_i(u) = \pi_{i,u}(q_i)_u^\top$, where $(q_i)_u^\top = (P_i[u, \cdot])_{-u}$. Thus, $\Delta_{i,u}(k_i) = -(\pi_{i,u})(q_i)_u^\top T_i(u)^{k_i-2} (p_i)_u \leq 0$ because any element in π_i , $(q_i)_u^\top$, $(p_i)_u$ and $P_i(u)$ is non-negative. $\square$

More interestingly, the stationary distribution π_i is the *only* starting distribution for P_i such that any $u \in \mathcal{V}_i$, $k_i \in \mathbb{Z}_{>0}$, $g_{i,u}(k_i + 1) - g_{i,u}(k_i) \leq 0$ when P_i is ergodic and there are no self loops in $\mathcal{G}_i$.

Proof. With a little abuse of the notation, we use P to denote the transition probability matrix $P_i \in \mathbb{R}^{n \times n}$ and let $P_{i,j}$ be the element in the i -th row and the j -th column. Since P is ergodic, according to Theorem 54 in Serfozo [34], there exists a unique and positive stationary distribution $\pi = (\pi_1, \dots, \pi_n)$, i.e., $\pi P = \pi$ and $\pi > 0$. Any starting distribution α can be represented by $\pi + \epsilon$, where $\epsilon = (\epsilon_1, \dots, \epsilon_n)$ is a perturbation vector, and $-\pi_j \leq \epsilon_j \leq 1 - \pi_j$, $j \in [n]$. We have the following equation for the margin of marginal gains for node u in the first two steps,

$$\begin{aligned} \Delta_u &= g_{i,u}(1) - g_{i,u}(2) = (\pi_u + \epsilon_u) - \sum_{j \neq u} (\pi_j + \epsilon_j) P_{j,u} \\ &= \pi_u P_{uu} + \epsilon_u - \sum_{j \neq u} \epsilon_j P_{j,u}. \end{aligned}$$

Since $\mathcal{G}$ has no self loops, i.e., $P_{u,u} = 0$, we have $\Delta_u = \epsilon_u - \sum_{j \neq u} \epsilon_j P_{j,u}$, and we can verify that $\sum_{u \in [n]} \Delta_u = \sum_{u \in [n]} (1 - \sum_{j \neq u} P_{u,j}) \epsilon_u = 0$. Hence, we have to guarantee $\Delta_u = 0$ for all u , otherwise there will exist a node u such that $\Delta_u < 0$. To ensure $\Delta_u = 0$, we need to ensure $\epsilon P = \epsilon$ by rephrasing the equations $\Delta_u = 0$ for all u . Again, since there exists a unique and positive stationary distribution for P and $\sum_{u \in [n]} \epsilon_u = 0$, we can derive $\epsilon = \beta \pi$, where β has to be 0. Therefore, combined with Lemma 2, the stationary distribution is the only starting distribution such that for any P_i , any $u \in \mathcal{V}_i$, $k_i \in \mathbb{Z}_{>0}$, $g_{i,u}(k_i + 1) - g_{i,u}(k_i) \leq 0$. $\square$

Appendix C. Proofs of Offline Optimization for Overlapping MuLaNE

C.1. Starting from the arbitrary distribution

Proof. By definition, we need to show $r_{\mathcal{G}, \alpha, \sigma}(x \wedge y) + r_{\mathcal{G}, \alpha, \sigma}(x \vee y) \leq r_{\mathcal{G}, \alpha, \sigma}(x) + r_{\mathcal{G}, \alpha, \sigma}(y)$ for any $x, y \in \mathbb{Z}_{\geq 0}^m$, and $r_{\mathcal{G}, \alpha, \sigma}(x) \leq r_{\mathcal{G}, \alpha, \sigma}(y)$ if $x \leq y$.

(Monotonicity.) By Eq. (7), $r_{\mathcal{G}, \alpha, \sigma}(k) = \sum_{v \in \mathcal{V}} \sigma_v (1 - \prod_{i \in [m]} (1 - P_{i,v}(k_i)))$. Since $P_{j,v}(x_j) \leq P_{j,v}(x_j + 1)$, for any $j \in [m]$, $v \in \mathcal{V}$, $x_j \in \mathbb{Z}_{\geq 0}$, we have

$$\begin{aligned} & r_{\mathcal{G}, \alpha, \sigma}(x + \chi_j) - r_{\mathcal{G}, \alpha, \sigma}(x) \\ &= \sum_{v \in \mathcal{V}} \sigma_v \left(\left(\prod_{i \neq j} (1 - P_{i,v}(x_i)) \right) (P_{j,v}(x_j + 1) - P_{j,v}(x_j)) \right) \\ &\geq 0, \end{aligned}$$

for any $x \in \mathbb{Z}_{\geq 0}^m$, $j \in [m]$. Then we can use the above inequality repeatedly to show that $r_{\mathcal{G}, \alpha, \sigma}(x) \leq r_{\mathcal{G}, \alpha, \sigma}(y)$ when $x \leq y$.

(Submodularity.) For submodular property, it suffices to prove that $(1 - \prod_{i \in [m]} (1 - P_{i,v}(k_i)))$ is submodular for any $v \in \mathcal{V}$, because a positive weighted sum of submodular function is still submodular. We will rely on the following lemma to prove the submodularity of $r_{\mathcal{G}, \alpha, \sigma}(k)$.

Lemma 5. Function $f : \mathbb{Z}_{\geq 0}^m \rightarrow \mathbb{R}$ is submodular if and only if

$$\begin{aligned} f(\mathbf{x} + \chi_i) - f(\mathbf{x}) &\geq f(\mathbf{x} + \chi_j + \chi_i) - f(\mathbf{x} + \chi_j), \\ \text{for any } \mathbf{x} \in \mathbb{Z}_{\geq 0}^m \text{ and } i \neq j. \end{aligned} \quad (\text{C.1})$$

Proof.

(If part.) We first prove, if Eq. (C.1) holds, the following inequality holds,

$$f(\mathbf{x} + \chi_i) - f(\mathbf{x}) \geq f(\mathbf{y} + \chi_i) - f(\mathbf{y}) \quad (\text{C.2})$$

for any $\mathbf{x} \leq \mathbf{y}$ and $i \in [m]$ such that $x_i = y_i$.

Let $I_0 = \{i \in [m] : x_i = y_i\}$, $I_1 = \{i \in [m] : x_i < y_i\}$. For any $\mathbf{x} \leq \mathbf{y}$, we denote the elements in I_1 by $i_1, \dots, i_s$ and write $\mathbf{y} = \mathbf{x} + \sum_{j=1}^s \alpha_j \chi_{i_j}$, where $\alpha_j = y_{i_j} - x_{i_j}$. For any $i \in I_0$, we have

$$\begin{aligned} f(\mathbf{x} + \chi_i) - f(\mathbf{x}) &\geq f(\mathbf{x} + \chi_{i_1} + \chi_i) - f(\mathbf{x} + \chi_{i_1}) \\ &\geq f(\mathbf{x} + 2\chi_{i_1} + \chi_i) - f(\mathbf{x} + 2\chi_{i_1}) \\ &\geq \dots \\ &\geq f(\mathbf{x} + \alpha_{i_1} \chi_{i_1} + \chi_i) - f(\mathbf{x} + \alpha_{i_1} \chi_{i_1}) \\ &\geq f(\mathbf{x} + \alpha_{i_1} \chi_{i_1} + \chi_{i_2} + \chi_i) - f(\mathbf{x} + \alpha_{i_1} \chi_{i_1} + \chi_{i_2}) \\ &\geq f(\mathbf{x} + \alpha_{i_1} \chi_{i_1} + \alpha_{i_2} \chi_{i_2} + \chi_i) - f(\mathbf{x} + \alpha_{i_1} \chi_{i_1} + \alpha_{i_2} \chi_{i_2}) \\ &\geq \dots \\ &\geq f(\mathbf{x} + \sum_{j=1}^s \alpha_j \chi_{i_j} + \chi_i) - f(\mathbf{x} + \sum_{j=1}^s \alpha_j \chi_{i_j}) \\ &= f(\mathbf{y} + \chi_i) - f(\mathbf{y}). \end{aligned} \quad (\text{C.3})$$

Then for any $i \in I_0$ and $a \in \mathbb{Z}_{\geq 0}$, we have

$$\begin{aligned} f(\mathbf{x} + a\chi_i) - f(\mathbf{x}) &= \sum_{j=1}^a (f(\mathbf{x} + j\chi_i) - f(\mathbf{x} + (j-1)\chi_i)) \\ &\geq \sum_{j=1}^a (f(\mathbf{y} + j\chi_i) - f(\mathbf{y} + (j-1)\chi_i)) \\ &\geq f(\mathbf{y} + a\chi_i) - f(\mathbf{y}). \end{aligned} \quad (\text{C.4})$$

The first inequality holds because of Eq. (C.3), the fact $\mathbf{x} + (j-1)\chi_i \leq \mathbf{y} + (j-1)\chi_i$ and $x_i + j-1 = y_i + j-1$.

Then for any $\mathbf{x}, \mathbf{y} \in \mathbb{Z}_{\geq 0}^m$, let $I_2 = \{i \in [m] : x_i > y_i\} = \{i_1, \dots, i_s\}$. We have

$$\begin{aligned} f(\mathbf{x}) - f(\mathbf{x} \wedge \mathbf{y}) &= \sum_{l=1}^s \left(f\left(\mathbf{x} \wedge \mathbf{y} + \sum_{j=1}^l (x_{i_j} - y_{i_j}) \chi_{i_j}\right) \right. \\ &\quad \left. - f\left(\mathbf{x} \wedge \mathbf{y} + \sum_{j=1}^{l-1} (x_{i_j} - y_{i_j}) \chi_{i_j}\right) \right) \\ &\geq \sum_{l=1}^s \left(f\left(\mathbf{y} + \sum_{j=1}^l (x_{i_j} - y_{i_j}) \chi_{i_j}\right) \right. \\ &\quad \left. - f\left(\mathbf{y} + \sum_{j=1}^{l-1} (x_{i_j} - y_{i_j}) \chi_{i_j}\right) \right) \\ &= f(\mathbf{x} \vee \mathbf{y}) - f(\mathbf{y}). \end{aligned}$$

The inequality is derived from Eq. (C.4) because $\mathbf{y} + \sum_{j=1}^{l-1} (x_{i_j} - y_{i_j}) \chi_{i_j} \geq \mathbf{x} \wedge \mathbf{y} + \sum_{j=1}^{l-1} (x_{i_j} - y_{i_j}) \chi_{i_j}$ for any $0 \leq l \leq s$ and $(\mathbf{x} \wedge \mathbf{y})_{i_l} = y_{i_l}$, which concludes the if part.

(Only if part.) Assume f is submodular, let $\mathbf{a} = \mathbf{x} + \chi_i, (\ominus)\mathbf{x} + \chi_j, i \neq j$, we have $\mathbf{a} \vee (\ominus)\mathbf{x} + \chi_i + \chi_j, \mathbf{a} \wedge (\ominus)\mathbf{x}$. $f(\mathbf{x} + \chi_i) - f(\mathbf{x}) = f(\mathbf{a}) - f(\mathbf{a} \wedge (\ominus)) \geq f(\mathbf{a} \vee (\ominus)) - f((\ominus)) = f(\mathbf{x} + \chi_j + \chi_i) - f(\mathbf{x} + \chi_j)$. $\square$

Then, by Lemma 5 and the explicit formula of the reward function given by Eq. (7), we can prove $g(\mathbf{x}, v) - g(\mathbf{x} + \chi_l, v) \geq g(\mathbf{x} + \chi_j, v) - g(\mathbf{x} + \chi_l + \chi_j, v)$ for any $\mathbf{x} \in \mathbb{Z}_{\geq 0}^m, v \in \mathcal{V}$ and $l \neq j \in [m]$, where $g(\mathbf{x}, v) = \prod_{i \in [m]} (1 - P_{i,v}(x_i))$. This holds due to the

fact that the left hand side equals to $\left(\prod_{i \in [m] \setminus \{j,l\}} (1 - P_{i,v}(x_i))\right)(1 - P_{j,v}(k_j))(P_{l,v}(k_l + 1) - P_{l,v}(k_l))$ and the right hand side equals to $\left(\prod_{i \in [m] \setminus \{j,l\}} (1 - P_{i,v}(x_i))\right)(1 - P_{j,v}(k_j + 1))(P_{l,v}(k_l + 1) - P_{l,v}(k_l))$, and the left hand side is larger or equal to the right hand side because $(1 - P_{j,v}(k_j)) \geq (1 - P_{j,v}(k_j + 1))$. By summation over all nodes $v \in \mathcal{V}$ with node weights $\sigma_v \in [0, 1]$, we can prove the reward function is submodular. $\square$

Proof. For theoretical analysis, we first give a modified version of Algorithm 1, which both provide the same solution $\mathbf{k}$ given the same problem instance $(\mathcal{G}, \alpha, \sigma, B, c)$. To see this fact, the modified Algorithm 1 considers invalid tentative allocations (i^*, b^*) (adding it will exceed the total budget constraint B) and remove them if $b^* > K$, while Algorithm 1 only considers valid allocations by directly removing invalid allocations in advance in line 10 of Algorithm 1.

With little abuse of the notation, we use $r(\mathbf{k})$ to represent the reward $r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k})$ of a given problem instance. Let $\mathbf{k}^* \in \mathbb{Z}_{\geq 0}^m$ denote the optimal budget allocation and $\mathbf{k}^j \in \mathbb{Z}_{\geq 0}^m$ denote the budget allocation before entering the j -th iteration of the while loop in the modified Algorithm 1. After entering the j -th iteration, the algorithm tries to extend the current budget allocation $\mathbf{k}^j$ by choosing the pair (i^*, b^*) , which we denote as (i^j, b^j) . Let s be the first iteration we can not extend the current solution, i.e., $\mathbf{k}^s = \mathbf{k}^{s+1}$ and $\mathbf{k}^j < \mathbf{k}^{j+1}$ for $j = 1, \dots, s-1$. If we can always extend the current solution, we set s to be $B+1$. For analysis, we temporarily add (i^s, b^s) (which is removed later on in the s -th iteration) to form a "virtual" budget allocation $\mathbf{k}^{s+1} = \mathbf{k}^s + b^s \chi_{i^s}$. Let $\mathbf{k}_g \in \mathbb{Z}_{\geq 0}^m$ denote the solution returned by the modified Algorithm 1.

We first introduce lemmas describing two important properties given by the submodularity over the integer lattice. $\square$

Lemma 6. Soma et al. [13]. Let $f : \mathbb{Z}_{\geq 0}^m \rightarrow \mathbb{R}$ be a submodular function. For any $\mathbf{x}, \mathbf{y} \in \mathbb{Z}_{\geq 0}^m$, we have,

$$\begin{aligned} f(\mathbf{x} \vee \mathbf{y}) \\ \leq f(\mathbf{x}) + \sum_{i \in \text{supp}^+(\mathbf{y}-\mathbf{x})} (f(\mathbf{x} + (y_i - x_i)\chi_i) - f(\mathbf{x})). \end{aligned} \quad (\text{C.5})$$

Lemma 7. Soma et al. [13]. Let $f : \mathbb{Z}_{\geq 0}^m \rightarrow \mathbb{R}$ be a monotone submodular function. For any $\mathbf{x}, \mathbf{y} \in \mathbb{Z}_{\geq 0}^m$ with $\mathbf{x} \leq \mathbf{y}$ and $i \in [m]$ we have,

$$f(\mathbf{x} \vee \mathbf{y}) - f(\mathbf{x}) \geq f(\mathbf{y} \vee \mathbf{x}) - f(\mathbf{y}). \quad (\text{C.6})$$

Then, we have the following lemma.

Lemma 8. For $j = 1, \dots, s$,

$$r(\mathbf{k}^{j+1}) \geq (1 - \frac{b^j}{B})r(\mathbf{k}^j) + \frac{b^j}{B}r(\mathbf{k}^*). \quad (\text{C.7})$$

Proof. This is because

$$\begin{aligned} r(\mathbf{k}^*) \\ \leq r(\mathbf{k}^* \vee \mathbf{k}^j) \\ \leq r(\mathbf{k}^j) + \sum_{i \in \text{supp}^+(\mathbf{k}^* - \mathbf{k}^j)} (r(\mathbf{k}^j + (k_i^* - k_i^j)\chi_i) - r(\mathbf{k}^j)) \\ = r(\mathbf{k}^j) + \sum_{i \in \text{supp}^+(\mathbf{k}^* - \mathbf{k}^j)} (r(\mathbf{k}^j + \alpha_i \chi_i) - r(\mathbf{k}^j)) \quad (\text{Let } \alpha_i = k_i^* - k_i^j) \\ \leq r(\mathbf{k}^j) + \sum_{i \in \text{supp}^+(\mathbf{k}^* - \mathbf{k}^j)} \left(\alpha_i \frac{r(\mathbf{k}^{j+1}) - r(\mathbf{k}^j)}{b^j} \right) \\ \leq r(\mathbf{k}^j) + B \left(\frac{r(\mathbf{k}^{j+1}) - r(\mathbf{k}^j)}{b^j} \right). \end{aligned} \quad (\text{C.8})$$

The second inequality comes from Lemma 6, the third inequality holds because of the greedy procedure, and the last inequality holds because $\sum_{i \in \text{supp}^+(\mathbf{k}^* - \mathbf{k}^j)} \alpha_i \leq B$. By rearranging terms, Eq. (C.7) holds. $\square$

Next, We can prove the following lemma.

Lemma 9. For $l = 1, \dots, s$,

$$\begin{aligned} r(\mathbf{k}^{l+1}) \geq r(\mathbf{k}^1) \prod_{j=1}^l \left(1 - \frac{b^j}{B} \right) \\ + r(\mathbf{k}^*) \left(1 - \prod_{j=1}^l \left(1 - \frac{b^j}{B} \right) \right). \end{aligned} \quad (\text{C.9})$$

Proof. We can prove this lemma by induction on l . When $l = 1$, the lemma holds due to Lemma 8. Assume that the lemma holds for $l-1$, we have the following inequality holds,

$$r(\mathbf{k}^{l+1})$$

$$\begin{aligned}
&\geq (1 - \frac{b^l}{B})r(\mathbf{k}^l) + \frac{b^l}{B}r(\mathbf{k}^*) \\
&\geq (1 - \frac{b^l}{B})r(\mathbf{k}^1) \prod_{j=1}^{l-1} \left(1 - \frac{b^j}{B}\right) \\
&\quad + (1 - \frac{b^l}{B})r(\mathbf{k}^*) \left(1 - \prod_{j=1}^{l-1} \left(1 - \frac{b^j}{B}\right)\right) + \frac{b^l}{B}r(\mathbf{k}^*) \\
&= r(\mathbf{k}^1) \prod_{j=1}^l \left(1 - \frac{b^j}{B}\right) + r(\mathbf{k}^*) \left(1 - \prod_{j=1}^l \left(1 - \frac{b^j}{B}\right)\right),
\end{aligned} \tag{C.10}$$

where the first inequality is due to Lemma 8 by setting $j = l$, and the second inequality is by the assumption for $l - 1$. By induction, Lemma 9 holds.

We then consider the following cases.

Case 1. Suppose the total budget used for $\mathbf{k}^s$ is larger or equal to ηB , i.e., $\sum_{j=1}^{s-1} b^j \geq \eta B$, where $\eta \in [0, 1]$.

We have the following inequality.

$$r(\mathbf{k}^s) \geq (1 - e^{-\eta})r(\mathbf{k}^*). \tag{C.11}$$

This is due to Lemma 9 by setting $l = s - 1$, combined with the fact that $\mathbf{k}^1 = \mathbf{0}$ and $\prod_{j=1}^{s-1} \left(1 - \frac{b^j}{B}\right) \leq e^{-\eta}$. The later fact holds because,

$$\begin{aligned}
\log \left(\prod_{j=1}^{s-1} \left(1 - \frac{b^j}{B}\right) \right) &= (s-1) \sum_{j=1}^{s-1} \frac{1}{s-1} \log \left(1 - \frac{b^j}{B}\right) \\
&\leq (s-1) \log \left(1 - \frac{1}{s-1} \sum_{j=1}^{s-1} \frac{b^j}{B}\right) \\
&\leq (s-1) \log \left(1 - \frac{\eta}{s-1}\right),
\end{aligned}$$

where the first inequality holds because of the Jensen's Inequality [35] and the second inequality holds because $\sum_{j=1}^{s-1} \frac{b^j}{B} \geq \eta$. Then we can easily check $\prod_{j=1}^{s-1} \left(1 - \frac{b^j}{B}\right) \leq (1 - \frac{\eta}{s-1})^{s-1} \leq e^{-\eta}$.

Case 2. Suppose the total budget $\sum_{j=1}^{s-1} b^j \leq \eta B$. Then, we have $b^s > (1 - \eta)B$. We can prove the following inequality holds,

$$r(\mathbf{k}_g) \geq (1 - \frac{1}{2 - \eta})r(\mathbf{k}^*). \tag{C.12}$$

This is due to Eq. (C.8), we have

$$\begin{aligned}
r(\mathbf{k}^*) &\leq r(\mathbf{k}^s) + B \left(\frac{r(\mathbf{k}^{s+1}) - r(\mathbf{k}^s)}{b^s} \right) \\
&\leq r(\mathbf{k}^s) + \left(\frac{r(\mathbf{k}^{s+1}) - r(\mathbf{k}^s)}{1 - \eta} \right) \\
&\leq (1 + \frac{1}{1 - \eta})r(\mathbf{k}_g),
\end{aligned}$$

where the second inequality is due to $b^s > (1 - \eta)B$ and the last equality is due to the fact $r(\mathbf{k}^{s+1}) - r(\mathbf{k}^s) \leq r(\mathbf{k}_g) - r(\mathbf{k}^s)$. To see the above fact, we assume without loss of generality the pair (i^s, b^s) improves $\mathbf{k}^s$ towards the optimal budget allocation, i.e., $k_{i^s}^{s+1} = b^s + k_{i^s}^s \leq k_{i^s}^*$. Otherwise, if $b^s + k_{i^s}^s > k_{i^s}^*$, we can safely delete (i^s, b^s) in the queue $\mathcal{Q}$ and does not affect the greedy solution, the optimal solution and the analysis. Let $r_{\max} = \max_{i \in [m]} r(c_i \chi_i)$, we have $r(\mathbf{k}^{s+1}) - r(\mathbf{k}^s) \leq r(\mathbf{k}^s \vee k_{i^s}^*, \chi_{i^s}) - r(\mathbf{k}^s) \leq r(k_{i^s}^* \chi_{i^s}) - r(\mathbf{0}) \leq r_{\max} \leq r(\mathbf{k}_g)$. Also, we can obtain that $r(\mathbf{k}^s) \leq r(\mathbf{k}_g)$ since $\mathbf{k}^s \leq \mathbf{k}_g$. By rearranging the terms, Eq. (C.12) holds.

Combining Eqs. (C.11) and (C.12), we have

$$\begin{aligned}
r(\mathbf{k}_g) &\geq \min_{\eta \in [0, 1]} \max \left(1 - e^{-\eta}, 1 - \frac{1}{2 - \eta} \right) r(\mathbf{k}^*) \\
&\geq (1 - e^{-\eta})r(\mathbf{k}^*),
\end{aligned}$$

where η is the solution for equation $e^\eta = 2 - \eta$. $\square$

C.2. Starting from the stationary distribution

Proof. By definition, we need to show $r_{\mathcal{G}, \pi, \sigma}(\mathbf{y} + \chi_j) - r_{\mathcal{G}, \pi, \sigma}(\mathbf{y}) \leq r_{\mathcal{G}, \pi, \sigma}(\mathbf{x} + \chi_j) - r_{\mathcal{G}, \pi, \sigma}(\mathbf{x})$, and $r_{\mathcal{G}, \pi, \sigma}(\mathbf{x}) \leq r_{\mathcal{G}, \pi, \sigma}(\mathbf{y})$ for any $\mathbf{x} \leq \mathbf{y}$, $j \in [m]$.

We first introduce the following lemma to help us to prove the DR-submodularity. $\square$

Lemma 10. Function f is DR-submodular if and only if

$$f(\mathbf{x} + \chi_i) - f(\mathbf{x}) \geq f(\mathbf{x} + \chi_j + \chi_i) - f(\mathbf{x} + \chi_j), \quad (\text{C.13})$$

for any $\mathbf{x} \in \mathbb{Z}_{\geq 0}^m$.

(If part.) We can easily check this direction holds by using the similar argument for Eq. (C.3), where the only difference is we consider $i \in [m]$ instead of $i \in I_0$.

(Only if part.) We can set $\mathbf{y} := \mathbf{x}' + \chi_j$, $\mathbf{x} := \mathbf{x}' + \chi_i$, for any $i, j \in [m]$, $\mathbf{x}' \in \mathbb{Z}_{\geq 0}^m$, and use Eq. (C.2) to show the only if part holds.

Since $r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k})$ is submodular for arbitrary starting distributions α , Eq. (C.1) holds. Then consider any layer $j \in [m]$ and budget allocation $\mathbf{x} \in \mathbb{Z}_{\geq 0}^m$, it is sufficient to show $r_{\mathcal{G}, \pi, \sigma}(\mathbf{x} + 2\chi_j) - r_{\mathcal{G}, \pi, \sigma}(\mathbf{x} + \chi_j) \leq r_{\mathcal{G}, \pi, \sigma}(\mathbf{x} + \chi_j) - r_{\mathcal{G}, \pi, \sigma}(\mathbf{x})$. Since the left hand side minus the right hand side equals to $\sum_{v \in \mathcal{V}} \sigma_v [(\prod_{i \in [m], i \neq j} (1 - P_{i,v}(x_i))) \cdot (P_{j,v}(x_j) - 2P_{j,u}(x_j + 1) + P_{j,v}(x_j + 2))]$, we only need to show $\sum_{v \in \mathcal{V}} \sigma_v [(\prod_{i \in [m], i \neq j} (1 - P_{i,v}(x_i))) \cdot (g_{j,v}(x_j + 2) - g_{j,v}(x_j + 1))] \leq 0$. The above inequality holds because of Lemma 2.

Proof. We can observe that the modified Algorithm 1 always select the pair (i^*, b^*) with $b^* = 1$ because $r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k} + b\chi_i) - r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k} + (b-1)\chi_i) \leq r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k} + \chi_i) - r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k})$ for arbitrary $\mathbf{k}$, i and b . Thus, we have $s = B + 1$ and by the similar argument for Eq. (C.11), we have $r(\mathbf{k}_g) = r(\mathbf{k}^{B+1}) \geq (1 - 1/e)r(\mathbf{k}^*)$, which completes the proof. $\square$

C.3. Efficiently evaluating the reward function

One key issue to derive the budget allocation is to efficiently evaluate the reward function $r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k})$ and its marginal gains $\delta(i, b, \mathbf{k})$. Since we need to repetitively use the visiting probabilities $P_{i,u}(k_i)$, we pre-calculate $P_{i,u}(k_i)$ in $O(m\|c\|_{\infty} n_{\max}^3)$ time in our algorithms based on Eq. (9), where $n_{\max} = \max_i |\mathcal{V}_i|$. For the overlapping case, given the current budget allocation $\mathbf{k}$, we maintain a value $p_u = \prod_{i \in [m]} (1 - P_{i,u}(k_i))$ for each node $u \in \mathcal{V}$. The marginal gain $\delta(i, b, \mathbf{k}) = \sum_{u \in \mathcal{V}} \sigma_u \left[p_u \left(1 - \frac{1 - P_{i,u}(k_i + b)}{1 - P_{i,u}(k_i)} \right) \right] / b$ can be evaluated in $O(n_{\max})$ time. Then, we update all $p_u = p_u \left(\frac{1 - P_{i,u}(k_i + b)}{1 - P_{i,u}(k_i)} \right)$ in $O(n_{\max})$ after we allocate b more budgets to layer i . Therefore, we can use $O(n_{\max})$ in total to evaluate $\delta(i, b, \mathbf{k})$. In practice, lazy evaluation [36] and parallel computing can be used to further accelerate our algorithm.

Appendix D. Budget effective greedy algorithm with partial enumeration and its analysis

D.1. Algorithm

The algorithm is shown in Algorithm 5.

Algorithm 5 Budget effective greedy algorithm with partial enumeration (BEGE).

Input: Graph $\mathcal{G}$, starting distributions α , node weights σ , budget B , constraints $\mathbf{c}$.

Output: Budget allocation $\mathbf{k}$.

- 1: Compute visiting probabilities $(P_{i,u}(b))_{i \in [m], u \in \mathcal{V}, b \in [c_i]}$ according to Eq. (7).
 - 2: $\mathbf{k} \leftarrow \text{BEGE}((P_{i,u}(b))_{i \in [m], u \in \mathcal{V}, b \in [c_i]}, \sigma, B, \mathbf{c})$.
 - 3: **Procedure** BEGE $((P_{i,u}(b))_{i \in [m], u \in \mathcal{V}, b \in [c_i]}, \sigma, B, \mathbf{c})$
 - 4: Let $\mathbf{k}_{\max} \leftarrow \mathbf{0}$.
 - 5: $S \leftarrow \{\mathbf{k} = (k_1, \dots, k_m) \mid 0 \leq k_i \leq c_i, \sum_{i \in [m]} k_i \leq B, \sum_{i \in [m]} \mathbb{I}\{k_i > 0\} \leq 3\}$. $\{S \text{ contains all partial solutions which allocate partial budgets to at most three layers}\}$
 - 6: **for** $\mathbf{k} \in S$ **do**
 - 7: $K \leftarrow B - \sum_{i \in [m]} k_i$.
 - 8: Let $Q \leftarrow \{(i, b_i) \mid i \in [m], 1 \leq b_i \leq c_i - k_i\}$.
 - 9: **while** $K > 0$ and $Q \neq \emptyset$ **do**
 - 10: $(i^*, b^*) \leftarrow \arg \max_{(i,b) \in Q} \delta(i, b, \mathbf{k}) / b$. {Eq. (11)}
 - 11: **if** $b^* \leq K$ **then**
 - 12: $k_{i^*} \leftarrow k_{i^*} + b^*$, $K \leftarrow K - b^*$.
 - 13: Modify all pairs $(i^*, b) \in Q$ to $(i^*, b - b^*)$.
 - 14: Remove all pairs $(i^*, b) \in Q$ such that $b \leq 0$.
 - 15: **else**
 - 16: Remove (i^*, b^*) from Q .
 - 17: **end if**
 - 18: **end while**
 - 19: **if** $r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k}) > r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k}_{\max})$, **then** $\mathbf{k}_{\max} \leftarrow \mathbf{k}$.
 - 20: **end for** **return** $\mathbf{k}_{\max} := (k_1, \dots, k_m)$.
 - 21: **end Procedure**
-

D.2. Analysis

Theorem 5. The Algorithm 5 obtains a $(1 - 1/e)$ -approximate solution to the overlapping MuLaNE with arbitrary starting distributions.

Proof. Suppose that we start from a partial solution $\mathbf{k}^1 \in \mathbb{Z}_{\geq 0}^m \neq \mathbf{0}$. Let us first reorder the optimal solution $\mathbf{k}^*$ according to a non-increasing marginal gain ordering. Namely, the marginal gain of the pair (s_1^*, b_1^*) with respect to the empty pair is the highest among all other pairs, the marginal gain of the pair (s_2^*, b_2^*) with respect to the solution consisting of pair (s_1^*, b_1^*) is the highest among all remaining pairs, and so on. To be concrete, $\mathbf{k}^* = \{(s_1^*, b_1^*), \dots, (s_m^*, b_m^*)\}$, where $s_i^* \in [m] \setminus \{s_1^*, \dots, s_{i-1}^*\}$ is selected to maximize the following equation:

$$\left(r\left(\sum_{j=1}^{i-1} b_j^* \chi_{s_j^*} + b_i^* \chi_{s_i^*}\right) - r\left(\sum_{j=1}^{i-1} b_j^* \chi_{s_j^*}\right) \right). \quad (\text{D.1})$$

Then, we try to bound the "virtual" marginal gain $\Delta^s = r(\mathbf{k}^{s+1}) - r(\mathbf{k}^s)$, and recall s is the first iteration we can not extend the current solution. Without loss of generality, we assume the virtual pair (i^s, b^s) improves $\mathbf{k}^s$ towards the optimal budget allocation, i.e., $b^s + k_{i^s}^s \leq k_{i^s}^*$. Otherwise, if $b^s + k_{i^s}^s > k_{i^s}^*$, we can safely delete (i^s, b^s) in the queue Q and does not affect the greedy solution, the optimal solution and the analysis.

If we start from the initial solution $\mathbf{k}^1$ such that $\mathbf{k}^1$ matches $(s_1^*, b_1^*), (s_2^*, b_2^*), (s_3^*, b_3^*)$, i.e., $k_{s_i^*}^1 = b_i^*$ for $i = 1, 2, 3$, we have $\Delta^s = r(\mathbf{k}^{s+1}) - r(\mathbf{k}^s) \leq r(\mathbf{k}^s \vee k_{i^s}^* \chi_{i^s}) - r(\mathbf{k}^s) \leq r(k_{i^s}^* \chi_{i^s}) - r(\mathbf{0}) \leq r(b_1^* \chi_{s_1^*})$. Moreover, $\Delta^s = r(\mathbf{k}^{s+1}) - r(\mathbf{k}^s) \leq r(\mathbf{k}^s \vee k_{i^s}^* \chi_{i^s}) - r(\mathbf{k}^s) = r(\mathbf{k}^s \vee b_1^* \chi_{s_1^*} \vee k_{i^s}^* \chi_{i^s}) - r(\mathbf{k}^s \vee b_1^* \chi_{s_1^*}) \leq r(b_1^* \chi_{s_1^*} \vee k_{i^s}^* \chi_{i^s}) - r(b_1^* \chi_{s_1^*}) \leq r(b_1^* \chi_{s_1^*} + b_2^* \chi_{s_2^*}) - r(b_1^* \chi_{s_1^*})$. Similarly, we have $\Delta^s \leq r(b_1^* \chi_{s_1^*} + b_2^* \chi_{s_2^*} + b_3^* \chi_{s_3^*}) - r(b_1^* \chi_{s_1^*} + b_2^* \chi_{s_2^*})$. By adding above inequalities, we have $\Delta^s \leq r(\mathbf{k}^1)/3$.

Now, we can use Lemma 9 by setting $l = s$ and the fact $\sum_{j=1}^s b^j \geq B$ to show $r(\mathbf{k}^{s+1}) \geq r(\mathbf{k}^1) + (1 - 1/e)(r(\mathbf{k}^*) - r(\mathbf{k}^1))$. Combining $\Delta^s \leq r(\mathbf{k}^1)/3$, we have $r(\mathbf{k}^s) \geq (1 - 1/3)r(\mathbf{k}^1) + (1 - 1/e)(r(\mathbf{k}^*) - r(\mathbf{k}^1)) \geq (1 - 1/e)r(\mathbf{k}^*)$, which completes the proof. $\square$

The time complexity of Algorithm 5 is $O(B^4 m^4 \|c\|_\infty n_{\max} + \|c\|_\infty m n_{\max}^3)$. This is because the number of all partial solutions is $|S| = O(B^3 m^3)$, and for any partial enumeration $\mathbf{k} \in S$, the time complexity is the same order as the Algorithm 1, i.e., $O(B\|c\|_\infty m n_{\max})$, where $O(n_{\max})$ is the time to evaluate $\delta(i, b, \mathbf{k})$ as discussed before.

Appendix E. Algorithms and analysis of offline optimization for non-overlapping MuLaNE

E.1. Starting from the stationary distribution

Algorithm 6 Myopic greedy algorithm for the non-overlapping MuLaNE.

Input: Network $\mathcal{G}$, starting distributions α , budget B , constraints c .

Output: Budget allocation $\mathbf{k}$.

- 1: Compute visiting probabilities $(P_{i,u}(b))_{i \in [m], u \in \mathcal{V}, b \in [c_i]}$ according to Eq. (8).
- 2: Compute layer-level marginal gain $g_i(b)_{i \in [m], b \in [c_i]}$ according to Eq. (E.2).
- 3: $\mathbf{k} \leftarrow \text{MG}((g_i(b))_{i \in [m], b \in [c_i]}, B, c)$

- 4: **Procedure** $\text{MG}((g_i(b))_{i \in [m], b \in [c_i]}, B, c)$
- 5: Let $\mathbf{k} := \{k_1, \dots, k_m\} \leftarrow \mathbf{0}$, $K \leftarrow B$.
- 6: **while** $K > 0$ **do**
- 7: $i^* \leftarrow \arg \max_{i \in [m], k_i + 1 \leq c_i} g_i(k_i + 1)$. $\{O(\log m)$ using the priority queue $\}$
- 8: $k_{i^*} \leftarrow k_{i^*} + 1$, $K \leftarrow K - 1$.
- 9: **end while** **return** $\mathbf{k} = (k_1, \dots, k_m)$.
- 10: **end Procedure**

According to Eqs. (8) and (10), the reward function can be rewritten as

$$r_{\mathcal{G}, \alpha, \sigma}(\mathbf{k}) = \sum_{i \in [m]} \sum_{b \in [k_i]} g_i(b), \quad (\text{E.1})$$

where $g_i(b)$ represents the layer-level marginal gain when we allocate one more budget (from $b - 1$ to b) to layer i , which is given by

$$g_i(b) = \sum_{v \in \mathcal{V}_i} \sigma_v(P_{i,v}(b) - P_{i,v}(b - 1)). \quad (\text{E.2})$$

From the definition of the reward function, we have two observations. First, budgets allocated to a layer will not affect the reward of other layers because layers are non-overlapping. Second, the layer-level marginal gain $g_i(b)$ for the i -th layer is non-increasing with respect to budget b . Based on these two observation, we propose the myopic greedy algorithm (Algorithm 6), which is a slight modification of Algorithm 2. Algorithm 6 takes the multi-layered network $\mathcal{G} = (\mathcal{G}_1, \dots, \mathcal{G}_m)$, starting distributions $(\alpha_1, \dots, \alpha_m)$, node weights $\sigma = (\sigma_1, \dots, \sigma_{|\mathcal{V}|})$ and total budget B as inputs. It consists of B rounds to compute the optimal budget allocation $\mathbf{k}^*$. In each round, line 7 in Algorithm 6 selects the layer i^* with the largest layer-level marginal gain and allocate one budget to layer i^* until total B budgets are used up.

Algorithm 7 Dynamic programming (DP) algorithm for the non-overlapping MuLaNE.

Input: Network $\mathcal{G}$, starting distributions α , node weights σ , budget B , constraints c .

Output: Budget allocation k .

```

1: Compute visiting probabilities ( $P_{i,u}(b)$ ) according to Eq. (8).
2: Compute layer-level marginal gain ( $g_i(b)$ ) according to Eq. (E.2).
3:  $k \leftarrow \text{DP}((g_i(b))_{i \in [m], b \in [c_i]}, B, c)$ 
4: Procedure  $\text{DP}((g_i(b))_{i \in [m], b \in [c_i]}, B, c)$ 
5: for  $i \in \{0\} \cup [m], b \in \{0\} \cup [B]$  do
6:   Set  $A[i, b] \leftarrow 0, V[i, b] \leftarrow 0$ .
7: end for
8: for  $i \leftarrow 0$  to  $m - 1$  do
9:   for  $b \in [B]$  do
10:     $j^* \leftarrow \arg \max_{j \in \{0\} \cup [c_{i+1}]} \{V[i, b - j] + r_{\mathcal{G}, \alpha, \sigma}(j \chi_{i+1})\}$ .
11:     $A[i + 1, b] \leftarrow A[i, b - j^*] + j^* \chi_{i+1}$ .
12:     $V[i + 1, b] \leftarrow V[i, b - j^*] + r_{\mathcal{G}, \alpha, \sigma}(j^* \chi_{i+1})$ .
13:   end for
14: end for
15: return  $A[m, B]$ .
16: end Procedure

```

Theorem 6. The Algorithm 6 obtains the optimal budget allocation to the non-overlapping MuLaNE with the stationary starting distributions.

Proof. Define a two dimensional array $\mathbf{M} \in \mathbb{R}^{c_1+c_2+\dots+c_m}$, where (i, j) -th entry is the j -th step layer-level marginal gain for layer i , i.e., $\mathbf{M}[i, j] = \sum_{v \in \mathcal{V}_i} g_{i,v}(j)$, for $i \in [m], j \in [c_i]$. Given the budget allocation $k = \{k_1, \dots, k_m\}$, the expected reward $r_{\mathcal{G}, \alpha, \sigma}(k)$ can be written as the sum of elements in $\mathbf{M}$, i.e., $r_{\mathcal{G}, \alpha, \sigma}(k) = \sum_{i=1}^m \sum_{j=1}^{k_i} \mathbf{M}[i, j]$. Because $g_{i,v}(k_i)$ is non-increasing with respect to k_i , the element in each row or the layer-level marginal gain is non-increasing. Hence, the greedy method at step s choose the s -th largest element in $\mathbf{M}$. At step $s = B$, the greedy policy selects all B largest elements and the corresponding budget allocation maximizes the expected reward $r_{\mathcal{G}, \alpha, \sigma}(k)$. $\square$

Algorithm 6 uses $O(\|c\|_{\infty} mn_{\max}^3)$ to compute visiting probabilities $P_{i,u}(k_i)$ and $O(\|c\|_{\infty} mn_{\max})$ to compute layer-wise marginal gains $g_i(b)$. Then $\delta(i, 1, k)$ can be evaluated in $O(\log m)$ using the priority queue, which is repeated for B iterations. Therefore, the time complexity of Algorithm 6 is $O(B \log m + \|c\|_{\infty} mn_{\max}^3)$.

E.2. Starting from arbitrary distributions

When random explorers start from any arbitrary distribution, Algorithm 6 can not obtain the optimal solution. This is because the layer-level marginal gain $g_i(b)$ is *not* non-increasing with respect to b . So we have to adopt a more general technique, dynamic programming (DP), to solve the budget allocation problem. The key idea is to keep a DP budget allocation table $A \in \mathbb{R}^{(m+1) \times (B+1)}$, where the (i, b) -th entry $A[i, b]$ saves the optimal budget allocation by allocating b budgets to the first i layers. Another DP table $V \in \mathbb{R}^{(m+1) \times (B+1)}$ saves the value of the optimal reward, where the (i, b) -th entry $V[i, b]$ corresponds to the optimal reward when setting $A[i, b]$ as the budget allocation. In the i' -th outer loop and b' -th inner loop in Algorithm 7, we have already obtained the optimal budget allocation for $A[i', b]$ with $i = 0, \dots, i', b \in [B]$, which help us to find the optimal amount of budget j^* to $(i' + 1)$ -th layer (in line 10), and thus we can obtain the $A[i' + 1, b']$ and $V[i' + 1, b']$ accordingly.

Theorem 7. Algorithm 7 obtains the optimal budget allocation to the non-overlapping MuLaNE with an arbitrary starting distribution.

Algorithm 7 uses $O(\|c\|_{\infty} mn_{\max}^3)$ to compute visiting probabilities $P_{i,u}(k_i)$ and $O(\|c\|_{\infty} mn_{\max})$ to compute layer-wise marginal gains $g_i(b)$. In the DP procedure, we can also pre-calculate all rewards $r_{\mathcal{G}, \alpha, \sigma}(j \chi_{i+1})$ in $O(\|c\|_{\infty} m)$. Then DP procedure takes $O(\|c\|_{\infty})$ to find j^* , which is repeated for Bm iterations. Therefore, the Algorithm 7 has the time complexity of $O(B\|c\|_{\infty} m + \|c\|_{\infty} mn_{\max}^3)$.

Appendix F. Analysis of online learning for overlapping MuLaNE

F.1. Proof of 1-Norm bounded smoothness

Proof. According to Eq. (7), and since $\mu_{i,u,b} \leq \mu'_{i,u,b}$, for $(i, u, b) \in \mathcal{A}$, we have $(1 - \prod_{i \in [m]} (1 - \mu_{i,u,k_i})) \leq (1 - \prod_{i \in [m]} (1 - \mu'_{i,u,k_i}))$. Therefore for any $\sigma_u \leq \sigma'_u$, we have $\sigma_u(1 - \prod_{i \in [m]} (1 - \mu_{i,u,k_i})) \leq \sigma'_u(1 - \prod_{i \in [m]} (1 - \mu'_{i,u,k_i}))$.

By summarizing both sides with $u \in \mathcal{V}$, we have $r_k(\mu, \sigma) \leq r_k(\mu', \sigma')$. $\square$

Proof. The left-hand side:

$$|r_{\mu, \sigma}(k) - r_{\mu', \sigma'}(k)|$$

$$\begin{aligned}
&= \left| \sum_{u \in \mathcal{V}} \left[\sigma_u \left(1 - \prod_{i \in [m]} (1 - \mu_{i,u,k_i}) \right) \right. \right. \\
&\quad \left. \left. - \sigma'_u \left(1 - \prod_{i \in [m]} (1 - \mu'_{i,u,k_i}) \right) \right] \right| \\
&\leq \sum_{u \in \mathcal{V}} \left| \sigma_u \left(1 - \prod_{i \in [m]} q_{i,u} \right) - \sigma'_u \left(1 - \prod_{i \in [m]} q'_{i,u} \right) \right| \\
&= \sum_{u \in \mathcal{V}} \left| \sigma_u \left(1 - \prod_{i \in [m]} q_{i,u} \right) - \sigma_u \left(1 - \prod_{i \in [m]} q'_{i,u} \right) \right. \\
&\quad \left. + \sigma_u \left(1 - \prod_{i \in [m]} q'_{i,u} \right) - \sigma'_u \left(1 - \prod_{i \in [m]} q'_{i,u} \right) \right| \\
&\leq \sum_{u \in \mathcal{V}} \sigma_u \left| \prod_{i \in [m]} q_{i,u} - \prod_{i \in [m]} q'_{i,u} \right| \\
&\quad + \sum_{u \in \mathcal{V}} |\sigma_u - \sigma'_u| \left(1 - \prod_{i \in [m]} q'_{i,u} \right),
\end{aligned}$$

where $q_{i,u} = (1 - \mu_{i,u,k_i})$, $q'_{i,u} = (1 - \mu'_{i,u,k_i})$.

For the first term, it can be derived that,

$$\begin{aligned}
&\sigma_u \left| \prod_{i \in [m]} q_{i,u} - \prod_{i \in [m]} q'_{i,u} \right| \\
&= \sigma_u \left| \sum_{i=1}^m \left(\left(\prod_{j=1}^{i-1} q'_{j,u} \right) (q_{i,u} - q'_{i,u}) \left(\prod_{j=i+1}^m q_{j,u} \right) \right) \right| \\
&\leq \sigma_u \sum_{i=1}^m \left| \left(\left(\prod_{j=1}^{i-1} q'_{j,u} \right) (q_{i,u} - q'_{i,u}) \left(\prod_{j=i+1}^m q_{j,u} \right) \right) \right| \\
&\leq \sigma_u \sum_{i=1}^m |q_{i,u} - q'_{i,u}| \\
&= \sigma_u \sum_{i \in [m]} |\mu_{i,u,k_i} - \mu'_{i,u,k_i}|,
\end{aligned}$$

where last inequality holds because $q_{j,u}, q'_{j,u} \leq 1$ and $\sigma_u \in [0, 1]$, for any $j \in [m], u \in \mathcal{V}$.

For the second term, we have

$$\begin{aligned}
&|\sigma_u - \sigma'_u| \left(1 - \prod_{i \in [m]} q'_{i,u} \right) \\
&= |\sigma_u - \sigma'_u| \left(1 - \prod_{i \in [m]} (1 - \mu'_{i,u,k_i}) \right) \\
&\leq |\sigma_u - \sigma'_u| \sum_{i \in [m]} \mu'_{i,u,k_i},
\end{aligned}$$

where the inequality holds because of the Weierstrass product inequality. Plugging the above two terms into the previous inequality, 1-Norm Bounded Smoothness is satisfied. $\square$

F.2. Regret proof analysis

In this section, we give the regret analysis for overlapping MuLaNE.

Let $\mathcal{A}$ denote the set containing all base arms, i.e., $\mathcal{A} = \{(i, u, b) | i \in [m], u \in \mathcal{V}, b \in [c_i]\}$. We add subscript t to denote the value of a variable at the end of round $t \in [T]$, e.g. $\bar{\sigma}_t, \hat{\mu}_{i,u,b,t}$, where T is the total number of rounds. For example, $T_{i,u,b,t}$ denotes the total number of times that arm (i, u, b) is played at the end of round t . Let us first introduce a definition of an unlikely event, where $\hat{\mu}_{i,u,b,t-1}$ is not accurate as expected.

Definition 2. We say that the sampling is nice at the beginning of round t , if for each arm $(i, u, b) \in \mathcal{A}$, $|\hat{\mu}_{i,u,b,t-1} - \mu_{i,u,b}| < \rho_{i,u,b,t}$, where $\rho_{i,u,b,t} = \sqrt{\frac{3 \ln t}{2T_{i,u,b,t-1}}}$. Let $\mathcal{N}_t$ be such event.

Lemma 11. For each round $t \geq 1$, $Pr\{\neg \mathcal{N}_t\} \leq 2|\mathcal{A}|t^{-2}$

Proof. For each round $t \geq 1$, we have

$$Pr\{\neg \mathcal{N}_t\}$$

$$\begin{aligned}
&= \Pr\{\exists(i, u, b) \in \mathcal{A}, |\hat{\mu}_{i,u,b,t-1} - \mu_{i,u,b}| \geq \sqrt{\frac{3 \ln t}{2T_{i,u,b,t-1}}}\} \\
&\leq \sum_{(i,u,b) \in \mathcal{A}} \Pr\{|\hat{\mu}_{i,u,b,t-1} - \mu_{i,u,b}| \geq \sqrt{\frac{3 \ln t}{2T_{i,u,b,t-1}}}\} \\
&= \sum_{(i,u,b) \in \mathcal{A}} \sum_{k=1}^{t-1} \Pr\{T_{i,u,b,t-1} = k \\
&\quad \wedge |\hat{\mu}_{i,u,b,t-1} - \mu_{i,u,b}| \geq \sqrt{\frac{3 \ln t}{2T_{i,u,b,t-1}}}\} \\
&\leq \sum_{(i,u,b) \in \mathcal{A}} \sum_{k=1}^{(t-1)} \frac{2}{t^3} \leq 2|\mathcal{A}|t^{-2}.
\end{aligned}$$

Given $T_{i,u,b,t-1} = k$, $\hat{\mu}_{i,u,b,t-1}$ is the average of k i.i.d. random variables $Y_{i,u,b}^{[1]}, \dots, Y_{i,u,b}^{[k]}$, where $Y_{i,u,b}^{[j]}$ is the Bernoulli random variable (computed in Algorithm 3 line 12) when the arm (i, u, b) is played for the j -th time during the execution. That is, $\hat{\mu}_{i,u,b,t-1} = \sum_{j=1}^k Y_{i,u,b}^{[j]}/k$. Note that the independence of $Y_{i,u,b}^{[1]}, \dots, Y_{i,u,b}^{[k]}$ comes from the fact we select a new starting position following the starting distribution α_t at the beginning of each round t . And the last inequality uses the Hoeffding Inequality [37]. $\square$

Then we can use monotonicity and 1-norm bounded smoothness properties (Properties 1 and 2) to bound the reward gap $\Delta_{k_t} = \xi r_{\mu,\sigma}(k^*) - r_{\mu,\sigma}(k_t)$ between the optimal action k^* and the action $k_t = (k_{t,1}, \dots, k_{t,m})$ selected by our algorithm $\mathcal{A}$. To achieve this, we introduce a positive real number $M_{i,u,b}$ for each arm (i, u, b) and define $M_{k_t} = \max_{(i,u,b) \in S_t} M_{i,u,b}$, where $S_t = \{(i, u, b) \in \mathcal{A} | b = k_{t,i}\}$. Define $|\mathcal{E}'| = |S_t| = \sum_{i \in [m]} |\mathcal{V}| = m|\mathcal{V}|$, and let

$$\kappa_T(M, s) = \begin{cases} 2, & \text{if } s = 0, \\ 3\sqrt{\frac{6 \ln T}{s}}, & \text{if } 1 \leq s \leq \ell_T(M), \\ 0, & \text{if } s \geq \ell_T(M) + 1, \end{cases}$$

where

$$\ell_T(M) = \lfloor \frac{54|\mathcal{E}'|^2 \ln T}{M^2} \rfloor.$$

Lemma 12. If $\{\Delta_{k_t} \geq M_{k_t}\}$ and $\mathcal{N}_t$ holds, we have

$$\Delta_{k_t} \leq \sum_{(i,u,b) \in S_t} \kappa_T(M_{i,u,b}, T_{i,u,b,t-1}). \quad (\text{F.1})$$

Proof. First, we can observe, for $(i, u, b) \in \mathcal{A}$, $\mu_{i,u,b}$ is upper bounded by $\bar{\mu}_{i,u,b,t}$, i.e., $\bar{\mu}_{i,u,b,t} \geq \mu_{i,u,b}$, when $\mathcal{N}_t$ holds. This is due to

$$\begin{aligned}
\bar{\mu}_{i,u,b,t} &= \max_{j \in [b]} \tilde{\mu}_{i,u,j,t} \\
&\geq \tilde{\mu}_{i,u,b,t} \\
&= \min\{\hat{\mu}_{i,u,b,t-1} + \rho_{i,u,b,t}, 1\} \\
&\geq \mu_{i,u,b},
\end{aligned}$$

where the last inequality comes from the fact that $|\hat{\mu}_{i,u,b,t-1} - \mu_{i,u,b}| < \rho_{i,u,b,t}$ when $\mathcal{N}_t$ holds.

Then, notice that the right hand of Eq. (F.1) is non-negative, it is trivially satisfied when $\Delta_{k_t} = 0$. We only need to consider $\Delta_{k_t} > 0$. By $\mathcal{N}_t$ and Property 1, We have the following inequalities,

$$r_{\bar{\mu},\bar{\sigma}}(k_t) \geq \xi r_{\bar{\mu},\bar{\sigma}}(k^*) \geq \xi r_{\mu,\sigma}(k^*) = \Delta_{k_t} + r_{\mu,\sigma}(k_t).$$

The first inequality is due to the oracle outputs the (ξ, β) -approximate solution given parameters $\bar{\mu}_t$, where $\xi = 1 - e^{-\eta}$, $\beta = 1$ by definition. The second inequality is due to Condition 1 (monotonicity).

Then, by Condition 2, we have

$$\begin{aligned}
\Delta_{k_t} &\leq r_{\bar{\mu},\bar{\sigma}}(k_t) - r_{\mu,\sigma}(k_t) \\
&\leq \sum_{(i,u,b) \in S_t} (\sigma_u |\bar{\mu}_{i,u,b,t} - \mu_{i,u,b}| + |\bar{\sigma}_u - \sigma_u| \bar{\mu}_{i,u,b,t}) \\
&\leq \sum_{(i,u,b) \in S_t} (|\bar{\mu}_{i,u,b,t} - \mu_{i,u,b}| + \rho_{i,u,b,t}),
\end{aligned}$$

where the last inequality is due to the following observation: $\bar{\sigma}_u \neq \sigma_u$ if and only if u has not been visited so far, so $\hat{\mu}_{i,u,b,t-1} = 0$ for any $i \in [m]$, $b \in [B]$ and $\bar{\mu}_{i,u,b,t} = \max_{j \in [b]} \{\hat{\mu}_{i,u,j,t-1} + \rho_{i,u,j,t}\} = \rho_{i,u,b,t}$.

Next, we can bound Δ_{k_t} by bounding $\bar{\mu}_{i,u,b,t} - \mu_{i,u,b} + \rho_{i,u,b,t}$ by applying a transformation. As $\{\Delta_{k_t} \geq M_{k_t}\}$ holds by assumption, so $\sum_{(i,u,b) \in S_t} |\bar{\mu}_{i,u,b,t} - \mu_{i,u,b}| + \rho_{i,u,b,t} \geq \Delta_{k_t} \geq M_{k_t}$. We have

$$\begin{aligned}
\Delta_{k_t} &\leq \sum_{(i,u,b) \in S_t} (|\bar{\mu}_{i,u,b,t} - \mu_{i,u,b}| + \rho_{i,u,b,t}) \\
&\leq -M_{k_t} + 2 \sum_{(i,u,b) \in S_t} |\bar{\mu}_{i,u,b,t} - \mu_{i,u,b}| + \rho_{i,u,b,t} \\
&= 2 \sum_{(i,u,b) \in S_t} (|\bar{\mu}_{i,u,b,t} - \mu_{i,u,b}| + \rho_{i,u,b,t} - \frac{M_{k_t}}{2|\mathcal{E}'|}) \\
&\leq 2 \sum_{(i,u,b) \in S_t} (|\bar{\mu}_{i,u,b,t} - \mu_{i,u,b}| + \rho_{i,u,b,t} - \frac{M_{i,u,b}}{2|\mathcal{E}'|}).
\end{aligned} \tag{F.2}$$

By $\mathcal{N}_t$, for all $(i, u, b) \in \mathcal{A}$, we have:

$$\begin{aligned}
&|\bar{\mu}_{i,u,b,t} - \mu_{i,u,b}| + \rho_{i,u,b,t} - \frac{M_{i,u,b}}{2|\mathcal{E}'|} \\
&= \tilde{\mu}_{i,u,l,t} - \mu_{i,u,b} + \rho_{i,u,b,t} - \frac{M_{i,u,b}}{2|\mathcal{E}'|} \quad (l = \arg \max_{j \in [b]} \tilde{\mu}_{i,u,j,t}) \\
&\leq \tilde{\mu}_{i,u,l,t} - \mu_{i,u,l} + \rho_{i,u,b,t} - \frac{M_{i,u,b}}{2|\mathcal{E}'|} \\
&\leq \min\{2\rho_{i,u,l,t} + \rho_{i,u,b,t}, 1\} - \frac{M_{i,u,b}}{2|\mathcal{E}'|} \\
&\leq \min\{3\rho_{i,u,b,t}, 1\} - \frac{M_{i,u,b}}{2|\mathcal{E}'|} \\
&\leq \min\left\{3\sqrt{\frac{3 \ln T}{2T_{i,u,b,t-1}}}, 1\right\} - \frac{M_{i,u,b}}{2|\mathcal{E}'|}.
\end{aligned}$$

where the first inequality is due to $\mu_{i,u,j} \leq \mu_{i,u,b}$ for $j \leq b$, the second inequality is due to the event $\mathcal{N}_t$ and the third inequality holds because $T_{i,u,j,t-1} \geq T_{i,u,b,t-1}$ for $j \leq b$.

Now consider the following cases:

Case 1. If $T_{i,u,b,t-1} \leq \ell_T(M_{i,u,b})$, we have $\bar{\mu}_{i,u,b,t} - \mu_{i,u,b} + \rho_{i,u,b,t} - \frac{M_{i,u,b}}{2|\mathcal{E}'|} \leq \min\left\{3\sqrt{\frac{3 \ln T}{2T_{i,u,b,t-1}}}, 1\right\} \leq \frac{1}{2}\kappa_T(M_{i,u,b}, T_{i,u,b,t-1})$.

Case 2. If $T_{i,u,b,t-1} \geq \ell_T(M_{i,u,b}) + 1$, then $3\sqrt{\frac{3 \ln T}{2T_{i,u,b,t-1}}} \leq \frac{M_{i,u,b}}{2|\mathcal{E}'|}$, so $\bar{\mu}_{i,u,b,t} - \mu_{i,u,b} + \rho_{i,u,b,t} - \frac{M_{i,u,b}}{2|\mathcal{E}'|} \leq 0 = \frac{1}{2}\kappa_T(M_{i,u,b}, T_{i,u,b,t-1})$.

Combine these two cases and Eq. (F.2), we have the following inequality, which concludes the lemma.

$$\Delta_{k_t} \leq \sum_{(i,u,b) \in S_t} \kappa_T(M_{i,u,b}, T_{i,u,b,t-1}).$$

□

Proof. By above lemmas, we have

$$\begin{aligned}
&\sum_{t=1}^T \mathbb{I}(\{\Delta_{k_t} \geq M_{k_t}\} \wedge \mathcal{N}_t) \Delta_{k_t} \\
&\leq \sum_{t=1}^T \sum_{(i,u,b) \in S_t} \kappa_T(M_{i,u,b}, T_{i,u,b,t-1}) \\
&= \sum_{(i,u,b) \in \mathcal{A}} \sum_{t' \in \{t | t \in [T], k_{t,j} = b\}} \kappa_T(M_{i,u,b}, T_{i,u,b,t'-1}) \\
&\leq \sum_{(i,u,b) \in \mathcal{A}} \sum_{s=0}^{T_{i,u,b,T}} \kappa_T(M_{i,u,b}, s) \\
&\leq 2|\mathcal{A}| + \sum_{(i,u,b) \in \mathcal{A}} \sum_{s=1}^{\ell_T(M_{i,u,b})} 3\sqrt{6\frac{\ln T}{s}} \\
&\leq 2|\mathcal{A}| + \sum_{(i,u,b) \in \mathcal{A}} \int_{s=0}^{\ell_T(M_{i,u,b})} 3\sqrt{6\frac{\ln T}{s}} ds \\
&\leq 2|\mathcal{A}| + \sum_{(i,u,b) \in \mathcal{A}} 6\sqrt{6 \ln T \ell_T(M_{i,u,b})}
\end{aligned}$$

$$\leq 2|\mathcal{A}| + \sum_{(i,u,b) \in \mathcal{A}} \frac{108|\mathcal{E}'| \ln T}{M_{i,u,b}}.$$

For distribution dependent bound, let us set $M_{i,u,b} = \Delta_{\min}^{i,u,b}$ and thus the event $\{\Delta_{k_t} \geq M_{k_t}\}$ always holds, i.e., $\mathbb{I}\{\Delta_{k_t} < M_{k_t}\} = 0$. We have

$$\begin{aligned} & \sum_{t=1}^T \mathbb{I}(\mathcal{N}_t \wedge \{\Delta_{k_t} \geq M_{k_t}\}) \Delta_{k_t} \\ & \leq 2|\mathcal{A}| + \sum_{(i,u,b) \in \mathcal{A}} \frac{108|\mathcal{E}'| \ln T}{\Delta_{\min}^{i,u,b}}. \end{aligned}$$

By Lemma 11, $Pr\{\neg \mathcal{N}_t\} \leq 2|\mathcal{A}|t^{-2}$, then

$$\mathbb{E} \left[\sum_{t=1}^T \mathbb{I}(\neg \mathcal{N}_t) \Delta_{k_t} \right] \leq \sum_{t=1}^T 2|\mathcal{A}|t^{-2} \Delta_{\max} \leq \frac{\pi^2}{3} |\mathcal{A}| \Delta_{\max}.$$

By our choice of $M_{i,u,b}$,

$$\sum_{t=1}^T \mathbb{I}(\Delta_{k_t} < M_{k_t}) \Delta_{k_t} = 0.$$

By the definition of the regret,

$$\begin{aligned} \text{Reg}_\mu(T) &= \mathbb{E} \left[\sum_{t=1}^T \Delta_{k_t} \right] \\ &\leq \mathbb{E} \left[\sum_{t=1}^T \left(\mathbb{I}(\neg \mathcal{N}_t) + \mathbb{I}(\Delta_{k_t} < M_{k_t}) \right) \right. \\ &\quad \left. + \mathbb{I}(\mathcal{N}_t \wedge \{\Delta_{k_t} \geq M_{k_t}\}) \Delta_{k_t} \right] \\ &\leq \sum_{(i,u,b) \in \mathcal{A}} \frac{108|\mathcal{E}'| \ln T}{\Delta_{\min}^{i,u,b}} + 2|\mathcal{A}| + \frac{\pi^2}{3} |\mathcal{A}| \Delta_{\max} \\ &= \sum_{(i,u,b) \in \mathcal{A}} \frac{108m|\mathcal{V}| \ln T}{\Delta_{\min}^{i,u,b}} + 2|\mathcal{A}| + \frac{\pi^2}{3} |\mathcal{A}| \Delta_{\max}. \end{aligned}$$

For distribution independent bound, we take $M_{i,u,b} = M = \sqrt{108|\mathcal{E}'||\mathcal{A}| \ln T / T}$, then we have $\sum_{t=1}^T \mathbb{I}\{\Delta_{k_t} \leq M_{k_t}\} \leq TM$. Then

$$\begin{aligned} \text{Reg}_\mu(T) &\leq \sum_{(i,u,b) \in \mathcal{A}} \frac{108|\mathcal{E}'| \ln T}{M_{i,u,b}} + 2|\mathcal{A}| + \frac{\pi^2}{3} |\mathcal{A}| \Delta_{\max} + TM \\ &\leq \frac{108|\mathcal{E}'||\mathcal{A}| \ln T}{M} + 2|\mathcal{A}| + \frac{\pi^2}{3} |\mathcal{A}| \Delta_{\max} + TM \\ &= 2\sqrt{108|\mathcal{E}'||\mathcal{A}| T \ln T} + \frac{\pi^2}{3} |\mathcal{A}| \Delta_{\max} + 2|\mathcal{A}| \\ &\leq 21\sqrt{m|\mathcal{V}||\mathcal{A}| T \ln T} + \frac{\pi^2}{3} |\mathcal{A}| \Delta_{\max} + 2|\mathcal{A}|. \quad \square \end{aligned}$$

F.3. Extension for random node weights

In this section, we extend the deterministic node weight σ_u for node $u \in \mathcal{V}$ to a random weight. Specifically, we denote $\bar{\sigma}_u \in [0, 1]$ as the random weight when we first visit u (after the first visit we will not get any reward so on so forth), and denote $\sigma_u := \mathbb{E}[\bar{\sigma}_u]$.

For the offline setting, the reward function in Eq. (7) remains the same since $\bar{\sigma}_u$ is independent of all other random variables. For the online setting, since the reward function remains unchanged, Properties 1 and 2 still hold. However, we need to change the way we maintain the optimistic node weight $\bar{\sigma}_u$. Compared with Algorithm 3, the major difference is that $\bar{\sigma}_u$ now represent the upper confidence bound (UCB) value of σ_u .

Regret analysis. Let $\mathcal{A}$ denote the set containing all base arms for visiting probabilities, i.e., $\mathcal{A} = \{(i, u, b) | i \in [m], u \in \mathcal{V}, b \in [c_i]\}$. With a bit of abuse of notation, we also use $\mathcal{V}$ to denote the set of base arms corresponding to the node weights. Therefore, $\mathcal{A} \cup \mathcal{V}$ is the set of all base arms, which is different from the Section F.2 that only has $\mathcal{A}$ as base arms and does not have base arms for random weights. Following the notations of Wang and Chen [17], we define $X_a \in [0, 1]$ be the random outcome of base arm $a \in \mathcal{A} \cup \mathcal{V}$ and

define D as the unknown joint distribution of random outcomes $X = (X_a)_{a \in \mathcal{A} \cup \mathcal{V}}$ with unknown mean $\mu_a = \mathbb{E}_{X \sim D}[X_a]$. Let k be any feasible budget allocation, we denote $p_a^{D,k}$ the probability that budget allocation k triggers arm a (i.e. outcome X_a is observed) when the unknown environment distribution is D . Specifically, $p_a^{D,k} = 1$ for $a \in \{(i, u, b) | i \in [m], u \in \mathcal{V}, b = k_i\}$, $p_a^{D,k} = 1 - \prod_{i \in [m]} (1 - \mu_{i,a,k_i})$ for $a \in \mathcal{V}$ and $p_a^{D,k} = 0$ for the rest of base arms. Therefore, we can rewrite the inequality in Property 2 as

$$\begin{aligned} & |r_{\mu,\sigma}(k) - r_{\mu',\sigma'}(k)| \\ & \leq \sum_{i \in [m], u \in \mathcal{V}, b = k_i} (\sigma'_u |\mu_{i,u,b} - \mu'_{i,u,b}|) \\ & + \sum_{u \in \mathcal{V}} |\sigma_u - \sigma'_u| \left(1 - \prod_{i \in [m]} (1 - \mu_{i,u,b}) \right) \\ & \leq \sum_{a \in \mathcal{A}} p_a^{D,k} |\mu_a - \mu'_a|, \end{aligned} \quad (\text{F.3})$$

where we abuse the notation to denote μ_a as σ_a and μ'_a as σ'_a if $a \in \mathcal{A}_2$ in the last inequality. We can observe that the above rewritten Property 2 can be viewed as a special case of the 1-norm TPM bounded smoothness condition in Wang and Chen [17].

Now we can apply the similar proof in Section B.3 in Wang and Chen [17], together with how we deal with the $|\bar{\mu}_{i,u,b,t} - \mu_{i,u,b}|$ in Appendix F.2 to get a distribution-dependent regret bound.

Let $K = m|\mathcal{V}| + |\mathcal{V}|$ be the maximum number of triggered arms, $|\mathcal{A} \cup \mathcal{V}| = Bm|\mathcal{V}| + |\mathcal{V}|$ be the number of base arms and $\lceil x \rceil_0 = \max\{\lceil x \rceil, 0\}$. For any feasible budget allocation k and any base arm $a \in \mathcal{A} \cup \mathcal{V}$, the gap $\Delta_k = \max\{0, \xi r_{\mu,\sigma}(k^*) - r_{\mu,\sigma}(k)\}$ and we define $\Delta_{\min}^a = \inf_{k: p_a^{D,k} > 0, \Delta_k > 0} \Delta_k$, $\Delta_{\max}^a = \sup_{k: p_a^{D,k} > 0, \Delta_k > 0} \Delta_k$. As convention, if there is no budget allocation k such that $p_a^{D,k} > 0$ and $\Delta_k > 0$, we define $\Delta_{\min}^a = +\infty$, $\Delta_{\max}^a = 0$. We define $\min_{a \in \mathcal{A} \cup \mathcal{V}} \Delta_{\min}^a$ and $\max_{a \in \mathcal{A} \cup \mathcal{V}} \Delta_{\max}^a$.

Theorem 8. Algorithm 3 with random node weights has the following distribution-dependent $(1 - e^{-\eta}, 1)$ approximation regret,

$$\begin{aligned} \text{Reg}_{\mu,\sigma}(T) & \leq \sum_{a \in \mathcal{A} \cup \mathcal{V}} \frac{576K \ln T}{\Delta_{\min}^a} \\ & + 4|\mathcal{A} \cup \mathcal{V}| + \sum_{a \in \mathcal{A} \cup \mathcal{V}} \left(\lceil \log_2 \frac{2K}{\Delta_{\min}^a} \rceil_0 + 2 \right) \frac{\pi^2}{6} \Delta_{\max}^a, \end{aligned}$$

We can also have a distribution-independent bound,

$$\begin{aligned} \text{Reg}_{\mu,\sigma}(T) & \leq 12\sqrt{|\mathcal{A} \cup \mathcal{V}|KT \ln T} \\ & + 2|\mathcal{A} \cup \mathcal{V}| + \left(\lceil \log_2 \frac{T}{18 \ln T} \rceil_0 + 2 \right) \frac{\pi^2}{6} \Delta_{\max}, \end{aligned}$$

Remark. Compared with the regret bound in Theorem 3 in the main text, we can see the base arms now become $\mathcal{A} \cup \mathcal{V}$ instead of $\mathcal{A}$, which is worse than the regret bound of the fixed unknown weight setting. This is a trade-off since we need more estimation to handle the uncertain random node weight, which incurs additional regrets. It also shows our non-trivial analysis of the fixed unknown weights carefully merges the error term into base arms $\mathcal{A}$ and incurs only a constant factor of extra regrets.

Appendix G. Analysis of explorer-relaxed online learning for overlapping MuLaNE

G.1. Proof of precision-balanced bounded smoothness

To prove this property while conserving space, we directly address a more general case, the MuLaNE problem with random node weights, rather than just fixed node weights.

Similar to Section F.3, let $\mathcal{A}$ denote the set containing all base arms for visiting probabilities, i.e., $\mathcal{A} = \{(i, u, b) | i \in [m], u \in \mathcal{V}, b \in [c_i]\}$. With a bit of abuse of notation, we also use $\mathcal{V}$ to denote the set of base arms corresponding to the node weights. Therefore, $\mathcal{A} \cup \mathcal{V}$ is the set of all base arms. Let k be any feasible budget allocation. We denote $p_a^{D,k}$ the probability that action k triggers arm a (i.e. outcome X_a is observed) when the unknown environment distribution is D . Specifically, $p_a^{D,k} = 1$ for $a \in \{(i, u, b) | i \in [m], u \in \mathcal{V}, b = k_i\}$, $p_a^{D,k} = 1 - \prod_{i \in [m]} (1 - \mu_{i,a,k_i})$ for $a \in \mathcal{V}$ and $p_a^{D,k} = 0$ for the rest of base arms.

Property 4. (General Precision-Balanced Bounded Smoothness). The reward function $r_{\mu,\sigma}(k)$ satisfies the following general precision-balanced bounded smoothness condition, i.e., for any budget allocation k , any two vectors $\mu = (\mu_{i,u,b})_{(i,u,b) \in \mathcal{A}}$, $\mu' = \mu + \zeta_\mu + \eta_\mu = (\mu'_{i,u,b})_{(i,u,b) \in \mathcal{A}}$ and any node weights $\sigma, \sigma' = \sigma + \zeta_\sigma + \eta_\sigma$, where $\zeta, \eta \in [-1, 1]^m$ represent the parameter change, we have $|r_{\mu,\sigma}(k) - r_{\mu',\sigma'}(k)| \leq$

$$\sum_{u \in \mathcal{V}} \left(\sum_{i \in [m], b = k_i} |\eta_{\mu,i,u,b}| + |\eta_{\sigma,u}| p_u^{D,k} \right) + \sqrt{1.25|\mathcal{V}|} \sqrt{\sum_{u \in \mathcal{V}} \left(\sum_{i \in [m], b = k_i} \left(\frac{\zeta_{\mu,i,u,b}^2}{(1 - \mu_{i,u,b})\mu_{i,u,b}} \right) + \frac{\zeta_{\sigma,u}^2 (p_u^{D,k})^2}{(1 - \sigma_u)\sigma_u} \right)}.$$

Unlike Condition 3 in Liu et al. [28] (Directional TPVM Bounded Smoothness), Property 4 represents the undirectional version (i.e., allowing ζ and η to be negative). We now provide the proof for Property 4. For fixed node weights, the same procedure applies by setting the triggering probability to 1 for the effective base arm.

Proof. Denote $\bar{p}_u^{D,k} = 1 - \prod_{i \in [m]} (1 - \mu'_{i,u,k_i})$ for $u \in \mathcal{V}$. To better-fit layout constraints and present the following complex proof clearly, we split the formula into two parts.

In the first part, we preliminary decompose and simplify the utility function difference:

$$\begin{aligned} & \left| r_{\mu, \sigma}(k) - r_{\mu', \sigma'}(k) \right| \\ &= \left| \sum_{u \in \mathcal{V}} \sigma'_u \left(\prod_{i \in [m]} (1 - \mu_{i, u, k_i}) - \prod_{i \in [m]} (1 - \mu'_{i, u, k_i}) \right) + \sum_{u \in \mathcal{V}} (\sigma'_u - \sigma_u) p_u^{D, k} \right| \end{aligned} \quad (\text{G.1})$$

$$\begin{aligned} &\leq \left| \sum_{u \in \mathcal{V}} \sum_{i \in [m]} (\mu'_{i, u, k_i} - \mu_{i, u, k_i}) \left(\prod_{\ell=1}^{i-1} (1 - \mu_{\ell, u, k_{\ell}}) \prod_{\ell=i+1}^m (1 - \mu_{\ell, u, k_{\ell}}) \right) \right| + \left| \sum_{u \in \mathcal{V}} (\sigma'_u - \sigma_u) p_u^{D, k} \right| \\ &\leq \sum_{u \in \mathcal{V}} \sum_{i \in [m], b=k_i} (|\zeta_{\mu, i, u, b}| + |\eta_{\mu, i, u, b}|) \prod_{\ell=1}^{i-1} (1 - \mu_{\ell, u, k_{\ell}}) \sigma'_u + \sum_{u \in \mathcal{V}} (|\zeta_{\sigma, u}| + |\eta_{\sigma, u}|) p_u^{D, k}, \end{aligned} \quad (\text{G.2})$$

where Eq. (G.1) is by telescoping the difference of $\prod_{i \in [m]} (1 - \mu_{i, u, k_i}) - \prod_{i \in [m]} (1 - \mu'_{i, u, k_i})$, and Eq. (G.2) is due to $\mu'_{i, u, k_i} \in [0, 1]$ for any $i \in [m], u \in \mathcal{V}$.

Then, we apply some common inequalities and deductions to bound the difference:

$$\begin{aligned} &\sum_{u \in \mathcal{V}} \sum_{i \in [m], b=k_i} (|\zeta_{\mu, i, u, b}| + |\eta_{\mu, i, u, b}|) \left(\prod_{\ell=1}^{i-1} (1 - \mu_{\ell, u, k_{\ell}}) \sigma'_u + \sum_{u \in \mathcal{V}} (|\zeta_{\sigma, u}| + |\eta_{\sigma, u}|) p_u^{D, k} \right) \\ &\leq \sum_{u \in \mathcal{V}} \sum_{i \in [m], b=k_i} \left(\frac{|\zeta_{\mu, i, u, b}|}{\sqrt{(1 - \mu_{i, u, b}) \mu_{i, u, b}}} \right) \sqrt{\left(\prod_{\ell=1}^{i-1} (1 - \mu_{\ell, u, k_{\ell}}) \mu_{i, u, b} \sigma'_u \right)} \\ &\quad + \sum_{u \in \mathcal{V}} \left(\frac{|\zeta_{\sigma, u}| p_u^{D, k}}{\sqrt{(1 - \sigma_u) \sigma_u}} \right) \sqrt{(1 - \sigma_u) \sigma_u} + \left(\sum_{u \in \mathcal{V}} \sum_{i \in [m], b=k_i} |\eta_{\mu, i, u, b}| + \sum_{u \in \mathcal{V}} |\eta_{\sigma, u}| p_u^{D, k} \right) \end{aligned} \quad (\text{G.3})$$

$$\begin{aligned} &\leq \sqrt{\sum_{u \in \mathcal{V}} \sum_{i \in [m], b=k_i} \left(\frac{\zeta_{\mu, i, u, b}^2}{(1 - \mu_{i, u, b}) \mu_{i, u, b}} \right) + \sum_{u \in \mathcal{V}} \frac{\zeta_{\sigma, u}^2 (p_u^{D, k})^2}{(1 - \sigma_u) \sigma_u}} \\ &\quad \cdot \sqrt{\sum_{u \in \mathcal{V}} \sum_{i \in [m], b=k_i} \prod_{\ell=1}^{i-1} (1 - \mu_{\ell, u, k_{\ell}}) \mu_{i, u, b} \sigma_u'^2 + \sum_{u \in \mathcal{V}} (1 - \sigma_u) \sigma_u} \\ &\quad + \left(\sum_{u \in \mathcal{V}} \sum_{i \in [m], b=k_i} |\eta_{\mu, i, u, b}| + \sum_{u \in \mathcal{V}} |\eta_{\sigma, u}| p_u^{D, k} \right). \end{aligned} \quad (\text{G.4})$$

In Eq. (G.3), we multiply the first term by $\sqrt{\mu_{i, u, b}}$ (where $b = k_i$) and then divide it by $\sqrt{\mu_{i, u, b}(1 - \mu_{i, u, b})}$. Similarly, we multiply the second term by $\sqrt{\sigma_u}$ and divide it by $\sqrt{\sigma_u(1 - \sigma_u)}$. In Eq. (G.4), we apply the Cauchy-Schwarz inequality simultaneously to both terms.

Then, with $\mu_{i, u, b} \in [0, 1], \forall i, u$ and $\sigma_u, \sigma'_u \in [0, 1], \forall u$, we have

$$\text{Eq. (G.4)} \quad (\text{G.5})$$

$$\begin{aligned} &\leq \sqrt{\sum_{u \in \mathcal{V}} \sum_{i \in [m], b=k_i} \left(\frac{\zeta_{\mu, i, u, b}^2}{(1 - \mu_{i, u, b}) \mu_{i, u, b}} \right) + \sum_{u \in \mathcal{V}} \frac{\zeta_{\sigma, u}^2 (p_u^{D, k})^2}{(1 - \sigma_u) \sigma_u}} \sqrt{\sum_{u \in \mathcal{V}} \sigma_u'^2 + (1 - \sigma_u) \sigma_u} \\ &\quad \left(\sum_{u \in \mathcal{V}} \sum_{i \in [m], b=k_i} |\eta_{\mu, i, u, b}| + \sum_{u \in \mathcal{V}} |\eta_{\sigma, u}| p_u^{D, k} \right) \end{aligned} \quad (\text{G.6})$$

$$\begin{aligned} &\leq \sqrt{\sum_{u \in \mathcal{V}} \sum_{i \in [m], b=k_i} \left(\frac{\zeta_{\mu, i, u, b}^2}{(1 - \mu_{i, u, b}) \mu_{i, u, b}} \right) + \sum_{u \in \mathcal{V}} \frac{\zeta_{\sigma, u}^2 (p_u^{D, k})^2}{(1 - \sigma_u) \sigma_u}} \cdot \sqrt{1.25 |\mathcal{V}|} \\ &\quad + \left(\sum_{u \in \mathcal{V}} \sum_{i \in [m], b=k_i} |\eta_{\mu, i, u, b}| + \sum_{u \in \mathcal{V}} |\eta_{\sigma, u}| p_u^{D, k} \right), \end{aligned} \quad (\text{G.7})$$

where Eq. (G.6) is due to the support of $\mu_{i, u, b}$ are bounded with $[0, 1]$ for any $i \in [m], u \in \mathcal{V}$ and Eq. (G.7) is due to the support of σ_u, σ'_u are bounded with $[0, 1]$ for any $u \in \mathcal{V}$. $\square$

G.2. Regret proof analysis

In this section, we provide detailed proofs for Theorem 4 under the more general case, i.e., random node weight, where we present an alternative, more precise proof for the main regret term compared to Liu et al. [28].

Let $\mathcal{A}$ denote the set that contains all base arms for visiting probabilities, i.e., $\mathcal{A} = \{(i, u, b) | i \in [m] \cup \mathcal{V}, u \in \mathcal{V}, b \in [c_i]\}$. With some abuse of notation, we also use $\mathcal{V}$ to denote the set of base arms corresponding to the node weights. Therefore, $\mathcal{A} \cup \mathcal{V}$ is the set of all base arms. Following the notations of Wang and Chen [17], we define $X_a \in [0, 1]$ be the random outcome of base arm $a \in \mathcal{A} \cup \mathcal{V}$ and define D as the unknown joint distribution of random outcomes $X = (X_a)_{a \in \mathcal{A} \cup \mathcal{V}}$ with unknown mean $\mu_a = \mathbb{E}_{X \sim D}[X_a]$. Let k be any feasible budget allocation, we denote $p_a^{D,k}$ the probability that budget allocation k triggers arm a (i.e. outcome X_a is observed) when the unknown environment distribution is D . Specifically, $p_a^{D,k} = 1$ for $a \in \{(i, u, b) | i \in [m], u \in \mathcal{V}, b = k_i\}$, $p_a^{D,k} = 1 - \prod_{i \in [m]} (1 - \mu_{i,a,k_i})$ for $a \in \mathcal{V}$ and $p_a^{D,k} = 0$ for the rest of base arms. Let $|\mathcal{A} \cup \mathcal{V}| = Bm|\mathcal{V}| + |\mathcal{V}|$ be the number of base arms and $\lceil x \rceil_0 = \max\{\lceil x \rceil, 0\}$. For any feasible budget allocation k and any base arm $a \in \mathcal{A} \cup \mathcal{V}$, the gap $\Delta_k = \max\{0, \xi r_{\mu,\sigma}(k^*) - r_{\mu,\sigma}(k)\}$ and we define $\Delta_{\min}^a = \inf_{k: p_a^{D,k} > 0, \Delta_k > 0} \Delta_k$, $\Delta_{\max}^a = \sup_{k: p_a^{D,k} > 0, \Delta_k > 0} \Delta_k$. By convention, if there is no action k such that $p_a^{D,k} > 0$ and $\Delta_k > 0$, we define $\Delta_{\min}^a = +\infty$, $\Delta_{\max}^a = 0$. We define $\Delta_{\min} = \min_{a \in \mathcal{A} \cup \mathcal{V}} \Delta_{\min}^a$ and $\Delta_{\max} = \max_{a \in \mathcal{A} \cup \mathcal{V}} \Delta_{\max}^a$. Denote $S = \{(i, u, b) \in \mathcal{A} | b = k_i\}$, and let $\tilde{S}$ be the set of base arms that can be triggered by S , i.e., $\{a \in \mathcal{A} \cup \mathcal{V} | p_a^{D,k} > 0\}$ under random node weights. The triggered arm set τ is extended to $\{(i, u, b) \in \mathcal{A} | k_i = b\} \cup \mathcal{V}$ here.

For fixed node weights, this can be considered a special case where $p_a^{D,k} = 1$ if $a \in \tilde{S}$, and $p_a^{D,k} = 0$ otherwise. Similarly to Section F.3, the set of base arms simplifies to $\mathcal{A} = \{(i, u, b) | i \in [m], u \in \mathcal{V}, b \in [c_i]\}$ under the fixed node weight setting. Consequently, due to the complexity of the proof for Theorem 4, which requires extensive detail, we will instead prove the more general and complex scenario with random node weights.

G.2.1. Useful concentration bounds, definitions and inequalities

We use the following tail bound for the construction of the confidence radius and our analysis.

Lemma 13 (Empirical Bernstein Inequality [38]). *Let $(X_a)_{a \in \mathcal{A} \cup \mathcal{V}}$ be n i.i.d random variables with bounded support $[0, 1]$ and mean $\mathbb{E}[X_a] = \mu$. Let $\hat{X}_n \triangleq \frac{1}{n} \sum_{a \in \mathcal{A} \cup \mathcal{V}} X_a$ and $\hat{Y}_n \triangleq \frac{1}{n} \sum_{i \in [n]} (X_a - \hat{X}_n)^2$ be the empirical mean and empirical variance of $(X_a)_{a \in \mathcal{A} \cup \mathcal{V}}$. Then for any $n \in \mathbb{N}$ and $y > 0$, it holds that*

$$\Pr \left[|\hat{X}_n - \mu| \geq \sqrt{\frac{2\hat{Y}_n y}{n}} + \frac{3y}{n} \right] \leq 3e^{-y} \quad (\text{G.8})$$

We use the following Bernstein Inequality to bound the difference between the empirical variance and the true variance.

Lemma 14 (Bernstein Inequality [37]). *Let $(X_a)_{a \in \mathcal{A} \cup \mathcal{V}}$ be n independent random variables in $[0, 1]$ with mean $\mathbb{E}[X_a] = \mu$ and variance $\text{Var}[X_a] \triangleq \mathbb{E}[X_a^2] - (\mathbb{E}[X_a])^2 = V$. Then with probability $1 - \delta$:*

$$\frac{1}{n} \sum_{a \in \mathcal{A} \cup \mathcal{V}} X_a \leq \mu + \frac{2 \log 1/\delta}{3n} + \sqrt{\frac{2V \log 1/\delta}{n}}. \quad (\text{G.9})$$

Following a similar approach to Wang and Chen [17], we define event-filtered regret to aid our analysis. However, since our proof does not require the definitions related to the triggering group analysis technique, we simplify the mathematical definitions accordingly.

Definition 3 (Event-Filtered Regret). For any series of events $(\mathcal{E}_t)_{t \geq [T]}$ indexed by round number t , we define the $\text{Reg}_{\xi, \mu, \sigma}(T, (\mathcal{E}_t)_{t \geq [T]})$ as the regret filtered by events $(\mathcal{E}_t)_{t \geq [T]}$, or the regret is only counted in t if $\mathcal{E}$ happens in t . Formally,

$$\text{Reg}_{\xi, \mu, \sigma}(T, (\mathcal{E}_t)_{t \geq [T]}) = \mathbb{E} \left[\sum_{t \in [T]} \mathbb{I}(\mathcal{E}_t) (\xi r_{\mu, \sigma}(k^*) - r_{\mu, \sigma}(k_t)) \right]. \quad (\text{G.10})$$

We will omit ξ, T and rewrite $\text{Reg}_{\xi, \mu, \sigma}(T, (\mathcal{E}_t)_{t \geq [T]})$ as $\text{Reg}_{\mu, \sigma}(T, \mathcal{E}_t)$ when contexts are clear for simplicity.

Definition 4. We say that the sampling is nice at the beginning of round t if: (1) for every base arm $a \in \mathcal{A} \cup \mathcal{V}$, $|\hat{\mu}_{a,t-1} - \mu_a| \leq \rho_{t,a}$, where $\rho_{t,a} = \sqrt{\frac{6\hat{Y}_{a,t-1} \log t}{T_{a,t-1}}} + \frac{9 \log t}{T_{a,t-1}}$; (2) for every base arm $a \in \mathcal{A} \cup \mathcal{V}$, $\hat{Y}_{a,t-1} \leq 2\mu_a(1 - \mu_a) + \frac{3.5 \log t}{T_{a,t-1}}$. We denote such event as $\mathcal{N}_t^s$.

The following lemma bounds the probability that $\mathcal{N}_t^s$ does not happen.

Lemma 15. For each round t , $\Pr[\neg \mathcal{N}_t^s] \leq 4|\mathcal{A} \cup \mathcal{V}|t^{-2}$.

Proof. Let $\mathcal{N}_t^{s,1}, \mathcal{N}_t^{s,2}$ be the event (1) and event (2), where $\mathcal{N}_t^s = \mathcal{N}_t^{s,1} \cap \mathcal{N}_t^{s,2}$. We first bound the probability that $\mathcal{N}_t^{s,1}$ does not happen, we have

$$\Pr[\neg \mathcal{N}_t^{s,1}] = \Pr \left[\exists a \in \mathcal{A} \cup \mathcal{V} \text{ s.t. } |\hat{\mu}_{a,t-1} - \mu_a| > \sqrt{\frac{6\hat{Y}_{a,t-1} \log t}{T_{a,t-1}}} + \frac{9 \log t}{T_{a,t-1}} \right] \quad (\text{G.11})$$

$$\leq \sum_{a \in \mathcal{A} \cup \mathcal{V}} \sum_{\tau \in [t]} \Pr \left[|\hat{\mu}_{a,t-1} - \mu_a| > \sqrt{\frac{6\hat{Y}_{a,t-1} \log t}{\tau}} + \frac{9 \log t}{\tau}, T_{a,t-1} = \tau \right] \quad (\text{G.12})$$

$$\leq 3|\mathcal{A} \cup \mathcal{V}|t^{-2}, \quad (\text{G.13})$$

where Eq. (G.12) is due to the union bound over i, τ , Eq. (G.13) is due to Lemma 13 by setting $y = 3 \log t$ and when $T_{a,t-1} = \tau$, $\hat{\mu}_{a,t-1}$ and $\hat{V}_{a,t-1}$ are the empirical mean and empirical variance of τ i.i.d random variables with mean μ_a .

We then bound the probability that the second event $\mathcal{N}_t^{s,2}$ does not happen using the similar proof of [39, Eq. (7)]. Fix $T_{a,t-1} = \tau$ and consider $(Y_a^1, \dots, Y_a^\tau)$, where $Y_a^k = (X_a^k - \mu_a)^2 \in [0, 1]$ and X_a^k is the random outcome of the k -th i.i.d trial. Since X_a^k are independent across k , Y_a^k are independent across k as well. In this case, one can verify that $\hat{V}_{a,t-1} = \frac{1}{\tau} \sum_{k=1}^{\tau} (X_a^k - \mu_a)^2 - (\frac{1}{\tau} \sum_{k=1}^{\tau} X_a^k - \mu_a)^2 \leq \frac{1}{\tau} \sum_{k=1}^{\tau} (X_a^k - \mu_a)^2 = \frac{1}{\tau} \sum_{k=1}^{\tau} \sigma_i^k$, $\mathbb{E}[Y_a^k] = \mathbb{E}[(X_a^k)^2] - \mu_a^2 \leq \mathbb{E}[X_a^k] \cdot 1 - \mu_a^2 = (1 - \mu_a)\mu_a$; and $\text{Var}[Y_a] = \mathbb{E}[(Y_a^k)^2] - (\mathbb{E}[Y_a^k])^2 \leq \mathbb{E}[(Y_a^k)^2] \leq \mathbb{E}[Y_a^k] \leq (1 - \mu_a)\mu_a$. By Lemma 14 over τ i.i.d random variable $(Y_a^k)_{k \in \tau}$, it holds with probability at least $1 - t^{-3}$ that

$$\frac{1}{\tau} \sum_{k=1}^{\tau} Y_a^k \leq \mathbb{E}[Y_a^k] + \frac{2 \log t}{\tau} + \sqrt{\frac{6 \text{Var}[Y_a^k] \log t}{\tau}} \quad (\text{G.14})$$

This implies

$$\hat{V}_{a,t-1} \leq \frac{1}{\tau} \sum_{k=1}^{\tau} Y_a^k \leq \mathbb{E}[Y_a^k] + \frac{2 \log t}{\tau} + \sqrt{\frac{6 \text{Var}[Y_a^k] \log t}{\tau}} \quad (\text{G.15})$$

$$\leq \mu_a(1 - \mu_a) + \frac{2 \log t}{\tau} + \sqrt{\frac{6(1 - \mu_a)\mu_a \log t}{\tau}} \quad (\text{G.16})$$

$$\leq \mu_a(1 - \mu_a) + \frac{2 \log t}{\tau} + \mu_a(1 - \mu_a) + \frac{3 \log t}{2\tau} \quad (\text{G.17})$$

$$= 2\mu_a(1 - \mu_a) + \frac{3.5 \log t}{\tau} \quad (\text{G.18})$$

where Eq. (G.16) is using $2xy \leq x^2 + y^2$ and $x = \sqrt{2\mu_a(1 - \mu_a)}$, $y = \sqrt{\frac{3 \log t}{n}}$.

Now by applying union bound over $a \in \mathcal{A} \cup \mathcal{V}$ and $\tau \in [t]$, we have $\Pr[\neg \mathcal{N}_t^{s,2}] \leq |\mathcal{A} \cup \mathcal{V}|t^{-2}$. Lastly, applying union bound over $\mathcal{N}_t^{s,1}$ and $\mathcal{N}_t^{s,2}$, we have $\Pr[\neg \mathcal{N}_t^s] \leq 4|\mathcal{A} \cup \mathcal{V}|t^{-2}$. $\square$

After setting up all above definitions, we can prove Lemma 16 about the confidence radius, which appears in the main content.

Lemma 16. Fix every base arm $(i, u, b) \in \mathcal{A} \cup \mathcal{V}$ and every time t , with probability at least $1 - 4|\mathcal{A} \cup \mathcal{V}|t^{-3}$, it holds that

$$\begin{aligned} \mu_{i,u,b} &\leq \min\{\hat{\mu}_{i,u,b,t-1} + \rho_{i,u,b,t}, 1\} =: \tilde{\mu}_{i,u,b,t} \leq \min\{\mu_{i,u,b} + 2\rho_{i,u,b,t}, 1\} \\ &\leq \min\left\{\mu_{i,u,b} + 4\sqrt{3}\sqrt{\frac{\mu_{i,u,b}(1 - \mu_{i,u,b}) \log t}{T_{i,u,b,t-1}}} + \frac{28 \log t}{T_{i,u,b,t-1}}, 1\right\}. \end{aligned} \quad (\text{G.19})$$

Proof. Recall that $\min\{\hat{\mu}_{i,u,b,t-1} + \rho_{i,u,b,t}, 1\} = \min\{\hat{\mu}_{i,u,b,t-1} + \sqrt{\frac{6\hat{V}_{i,u,b,t-1} \log t}{T_{i,u,b,t-1}}} + \frac{9 \log t}{T_{i,u,b,t-1}}, 1\}$. Under event $\mathcal{N}_t^{s,1}$, we have $|\mu_{i,u,b} - \hat{\mu}_{i,u,b,t}| \leq \rho_{i,u,b,t}$ by Lemma 15, hence the first and second inequality in Lemma 16 holds.

For the last inequality, under event $\mathcal{N}_t^{s,2}$, it holds that

$$\mu_{i,u,b} + 2\rho_{i,u,b,t} = \mu_{i,u,b} + 2\left(\sqrt{\frac{6\hat{V}_{i,u,b,t-1} \log t}{T_{i,u,b,t-1}}} + \frac{9 \log t}{T_{i,u,b,t-1}}\right) \quad (\text{G.20})$$

$$\leq \mu_{i,u,b} + 2\left(\sqrt{\frac{6 \cdot (2\mu_{i,u,b}(1 - \mu_{i,u,b}) + \frac{3.5 \log t}{T_{i,u,b,t-1}}) \log t}{T_{i,u,b,t-1}}} + \frac{9 \log t}{T_{i,u,b,t-1}}\right) \quad (\text{G.21})$$

$$\leq \mu_{i,u,b} + 4\sqrt{3}\sqrt{\frac{\mu_{i,u,b}(1 - \mu_{i,u,b}) \log t}{T_{i,u,b,t-1}}} + 2\sqrt{21}\frac{\log t}{T_{i,u,b,t-1}} + \frac{18 \log t}{T_{i,u,b,t-1}} \quad (\text{G.22})$$

$$\leq \mu_{i,u,b,t-1} + 4\sqrt{3}\sqrt{\frac{\mu_{i,u,b}(1 - \mu_{i,u,b}) \log t}{T_{i,u,b,t-1}}} + \frac{28 \log t}{T_{i,u,b,t-1}}, \quad (\text{G.23})$$

where Eq. (G.22) uses $\sqrt{x+y} \leq \sqrt{x} + \sqrt{y}$.

Since $\mathcal{N}_t^s = \mathcal{N}_t^{s,1} \cap \mathcal{N}_t^{s,2}$ and by Lemma 15, Eq. (G.19) holds with probability at least $1 - 4|\mathcal{A} \cup \mathcal{V}|t^{-2}$. $\square$

G.2.2. Decompose the total regret to event-filtered regrets

We firstly decompose the regret $\text{Reg}_{\mu,\sigma}(T, \{\}) = \text{Reg}_{\mu,\sigma}(T, \mathcal{N}_t^s, \mathcal{N}_t^o) + \text{Reg}_{\mu,\sigma}(T, \neg(\mathcal{N}_t^s \cap \mathcal{N}_t^o)) \leq \text{Reg}_{\mu,\sigma}(T, \mathcal{N}_t^s, \mathcal{N}_t^o) + \text{Reg}_{\mu,\sigma}(T, \neg \mathcal{N}_t^s) + \text{Reg}_{\mu,\sigma}(T, \neg \mathcal{N}_t^o)$, where $\mathcal{N}_t^s$ is defined in Lemma 4, $\mathcal{N}_t^o$ denotes the event where oracle successfully outputs an ξ -approximate solution (with probability at least β). We have the following lemma to do the decomposition.

Lemma 17 (Leading Regret Term). *Define the error term*

$$e_t(k_t) = 4\sqrt{3.75|\mathcal{V}|} \sqrt{\sum_{a \in \tilde{S}_t} (\frac{\log t}{T_{a,t-1}} \wedge \frac{1}{28})(p_a^{D,k_t})^2} + 28 \sum_{a \in \tilde{S}_t} (\frac{\log t}{T_{a,t-1}} \wedge \frac{1}{28})(p_a^{D,S_t}) \quad (\text{G.24})$$

and event $E_t = \{\Delta_{k_t} \leq e_t(k_t)\}$. The regret of Algorithm 4, when used with $(1 - e^{-\eta}, 1)$ approximation oracle is bounded by

$$\text{Reg}_{\mu,\sigma}(T) \leq \text{Reg}_{\mu,\sigma}(T, E_t) + \frac{2\pi^2}{3} |\mathcal{A} \cup \mathcal{V}| \Delta_{\max}. \quad (\text{G.25})$$

Proof. Recall that $k_t = (k_{t,1}, \dots, k_{t,m})$, $S_t = \{(i, u, b) \in \mathcal{A} \mid b = k_{t,i}\}$, and $\tilde{S}_t$ represents the base arms that can be triggered by S_t under the random node weight setting. Under event $\mathcal{N}_t^s, \mathcal{N}_t^o$, by Lemma 16, it is easily to check that

$$\begin{aligned} \min\{\mu_{a,t-1} + \rho_{a,t}, 1\} &\leq \bar{\mu}_{a,t-1} \\ &\leq \min\{\mu_{a,t-1} + 4\sqrt{3} \sqrt{\frac{\mu_a(1-\mu_a)\log t}{T_{a,t-1}}} + \frac{28\log t}{T_{a,t-1}}, 1\} \\ &\leq \mu_{a,t-1} + 4\sqrt{3} \sqrt{\mu_a(1-\mu_a)(\frac{\log t}{T_{a,t-1}} \wedge \frac{1}{28})} \\ &\quad + 28(\frac{\log t}{T_{a,t-1}} \wedge \frac{1}{28}) \end{aligned} \quad (\text{G.26})$$

Therefore, it holds that

$$\xi r_{\mu,\sigma}(k^*) \leq \xi r_{\bar{\mu},\bar{\sigma}}(k^*) \leq r_{\bar{\mu},\bar{\sigma}}(k_t) \quad (\text{G.27})$$

$$\begin{aligned} &\leq r_{\mu_t,\sigma}(k_t) + 4\sqrt{3.75|\mathcal{V}|} \sqrt{\sum_{a \in \tilde{S}_t} (\frac{\log t}{T_{a,t-1}} \wedge \frac{1}{28})(p_a^{D,k_t})^2} \\ &\quad + 28 \sum_{a \in \tilde{S}_t} (\frac{\log t}{T_{a,t-1}} \wedge \frac{1}{28}) p_a^{D,k_t}, \end{aligned} \quad (\text{G.28})$$

where the first inequality in Eq. (G.27) is due to monotonicity (Property 1), and second inequality in Eq. (G.27) is due to event $\mathcal{N}_t^o$, Eq. (G.28) is because of Eq. (G.26) and Property 4 by plugging in $\zeta_a = 4\sqrt{3} \sqrt{\mu_a(1-\mu_a)(\frac{\log t}{T_{a,t-1}} \wedge \frac{1}{28})}$ and $\eta_a = 28(\frac{\log t}{T_{a,t-1}} \wedge \frac{1}{28})$.

So $\text{Reg}_{\mu,\sigma}(T, \mathcal{N}_t^s, \mathcal{N}_t^o) \leq \text{Reg}_{\mu,\sigma}(T, E_t)$. Now for $\text{Reg}_{\mu,\sigma}(T, \neg \mathcal{N}_t^s)$, by Lemma 15 it holds that

$$\text{Reg}_{\mu,\sigma}(T, \neg \mathcal{N}_t^s) \leq \sum_{t=1}^T \Pr[\neg \mathcal{N}_t^s] \leq \sum_{t=1}^T 4|\mathcal{A} \cup \mathcal{V}| t^{-2} \leq \frac{2\pi^2}{3} |\mathcal{A} \cup \mathcal{V}| \Delta_{\max}. \quad (\text{G.29})$$

Similarly by definition, it holds that

$$\text{Reg}_{\mu,\sigma}(T, \neg \mathcal{N}_t^o) \leq (1 - \beta)T \Delta_{\max}. \quad (\text{G.30})$$

Therefore $\text{Reg}_{\mu,\sigma}(T, \{\}) \leq \text{Reg}_{\mu,\sigma}(T, E_t) + \frac{2\pi^2}{3} |\mathcal{A} \cup \mathcal{V}| \Delta_{\max} + (1 - \beta)T \Delta_{\max}$. And we have $\text{Reg}_{\mu,\sigma}(T) = \text{Reg}_{\mu,\sigma}(T, \{\}) - (1 - \beta)T \Delta_{\max} \leq \text{Reg}_{\mu,\sigma}(T, E_t) + \frac{2\pi^2}{3} \Delta_{\max}$, which concludes Lemma 17. $\square$

Recall event $E_t = \{\Delta_{k_t} \leq e_t(k_t)\}$, where $e_t(k_t) = 28 \sum_{a \in \tilde{S}_t} (\frac{\log t}{T_{a,t-1}} \wedge \frac{1}{28})(p_a^{D,k_t}) + 4\sqrt{3.75|\mathcal{V}|} \sqrt{\sum_{a \in \tilde{S}_t} (\frac{\log t}{T_{a,t-1}} \wedge \frac{1}{28})(p_a^{D,k_t})^2}$. We will further decompose the event-filtered regret $\text{Reg}_{\mu,\sigma}(T, E_t)$ into two event-filtered regret $\text{Reg}_{\mu,\sigma}(T, E_{t,1})$ and $\text{Reg}_{\mu,\sigma}(T, E_{t,2})$,

$$\text{Reg}_{\mu,\sigma}(T, E_t) \leq \text{Reg}_{\mu,\sigma}(T, E_{t,1}) + \text{Reg}_{\mu,\sigma}(T, E_{t,2}), \quad (\text{G.31})$$

where event $E_{t,1} = \{\Delta_{k_t} \leq 2e_{t,1}(k_t)\}$, event $E_{t,2} = \{\Delta_{k_t} \leq 2e_{t,2}(k_t)\}$, $e_{t,1}(k_t) = 4\sqrt{3.75|\mathcal{V}|} \sqrt{\sum_{a \in \tilde{S}_t} (\frac{\log t}{T_{a,t-1}} \wedge \frac{1}{28})(p_a^{D,k_t})^2}$, $e_{t,2}(k_t) = 28 \sum_{a \in \tilde{S}_t} (\frac{\log t}{T_{a,t-1}} \wedge \frac{1}{28})(p_a^{D,k_t})$. The above inequality holds since the following facts: We can observe $e_{t,1}(k_t) + e_{t,2}(k_t) = e_t(k_t)$. From E_t , we know either $E_{t,1}$ holds or $E_{t,2}$ holds. So E_t implies $1 \leq \mathbb{I}\{E_{t,1}\} + \mathbb{I}\{E_{t,2}\}$, and thus $\Delta_{k_t} \mathbb{I}\{E_t\} \leq \Delta_{k_t} \mathbb{I}\{E_{t,1}\} + \Delta_{k_t} \mathbb{I}\{E_{t,2}\}$, which concludes $\text{Reg}_{\mu,\sigma}(T, E_t) \leq \text{Reg}_{\mu,\sigma}(T, E_{t,1}) + \text{Reg}_{\mu,\sigma}(T, E_{t,2})$. The next section will provide two proofs for $\text{Reg}_{\mu,\sigma}(T, E_{t,1})$, and $\text{Reg}_{\mu,\sigma}(T, E_{t,2})$, respectively.

G.3. Refined regret bound analysis

In this section, we are going to bound the $\text{Reg}_{\mu,\sigma}(T, E_{t,1})$ and $\text{Reg}_{\mu,\sigma}(T, E_{t,2})$ separately under the event $\mathcal{N}_t^t$. The approach leverages a refined reverse amortization technique from Wang and Chen [17], Liu et al. [28] to assign regret Δ_{k_t} to each base arm based on carefully designed thresholds. Deriving the correct thresholds and regret allocation strategy is highly non-trivial, as it aims to minimize the factors m , $|\mathcal{V}|$, and T for the MuLaNE problem. In contrast to Liu et al. [28], we have eliminated a logarithmic factor related to the number of explorers from the second main term in their regret upper bound. Finally, we offer discussions on the distribution-independent regret bounds.

G.3.1. Bounding the regret on the first term

Let $c_1 = 4\sqrt{3}$, and $\tilde{\Delta}_{k_t} = \Delta_{k_t}/p_a^{D,k_t}$. Given filtration $\mathcal{F}_{t-1}$ and under event $E_{t,1}$, we have

$$\Delta_{k_t} \leq \sum_{a \in \mathcal{A} \cup \mathcal{V}} \frac{5c_1^2 |\mathcal{V}| (p_a^{D,k_t})^2 \frac{\log t}{T_{a,t-1}}}{\Delta_{k_t}} \quad (\text{G.32})$$

$$= -\Delta_{k_t} + 2 \sum_{a \in \mathcal{A} \cup \mathcal{V}} \frac{5c_1^2 |\mathcal{V}| (p_a^{D,k_t})^2 \frac{\log t}{T_{a,t-1}}}{\Delta_{k_t}} \quad (\text{G.33})$$

$$= -\frac{\sum_{a \in \tilde{\mathcal{S}}_t} p_a^{D,k_t} \Delta_{k_t} / p_a^{D,k_t}}{m|\mathcal{V}| + |\mathcal{V}|} + 2 \sum_{a \in \mathcal{A} \cup \mathcal{V}} \frac{5c_1^2 |\mathcal{V}| (p_a^{D,k_t})^2 \frac{\log t}{T_{a,t-1}}}{\Delta_{k_t}} \quad (\text{G.34})$$

$$\leq \sum_{a \in \mathcal{A} \cup \mathcal{V}} p_a^{D,k_t} \left(\frac{10c_1^2 |\mathcal{V}| \frac{\log t}{T_{a,t-1}}}{\Delta_{k_t} / p_a^{D,k_t}} - \frac{\Delta_{k_t} / p_a^{D,k_t}}{m|\mathcal{V}| + |\mathcal{V}|} \right) \quad (\text{G.35})$$

$$= \sum_{a \in \mathcal{A} \cup \mathcal{V}} p_a^{D,k_t} \left(\frac{10c_1^2 |\mathcal{V}| \frac{\log t}{T_{a,t-1}}}{\tilde{\Delta}_{k_t}} - \frac{\tilde{\Delta}_{k_t}}{m|\mathcal{V}| + |\mathcal{V}|} \right). \quad (\text{G.36})$$

$$\Delta_{k_t} = \mathbb{E}_t[\Delta_{k_t}] \leq \mathbb{E}_t \left[\sum_{a \in \mathcal{A} \cup \mathcal{V}} p_a^{D,k_t} \left(\frac{10c_1^2 |\mathcal{V}| \frac{\log t}{T_{a,t-1}}}{\tilde{\Delta}_{k_t}} - \frac{\tilde{\Delta}_{k_t}}{m|\mathcal{V}| + |\mathcal{V}|} \right) \right] \quad (\text{G.37})$$

$$= \mathbb{E}_t \left[\sum_{a \in \tau_t} \left(\frac{10c_1^2 |\mathcal{V}| \frac{\log t}{T_{a,t-1}}}{\tilde{\Delta}_{k_t}} - \frac{\tilde{\Delta}_{k_t}}{m|\mathcal{V}| + |\mathcal{V}|} \right) \right] \quad (\text{G.38})$$

$$\leq \mathbb{E}_t \left[\sum_{a \in \tau_t} \kappa_{a,T}(T_{a,t-1}) \right]. \quad (\text{G.39})$$

Regarding the last inequality, we define a regret allocation function

$$\kappa_{a,T}(\ell) = \begin{cases} \frac{1.25c_1^2 |\mathcal{V}|}{\tilde{\Delta}_{\min}^a}, & \text{if } \ell = 0, \\ 2\sqrt{\frac{5c_1^2 |\mathcal{V}| \log T}{\ell}}, & \text{if } 1 \leq \ell \leq L_{a,T,1}, \\ \frac{10c_1^2 |\mathcal{V}| \log T}{\tilde{\Delta}_{\min}^a \ell}, & \text{if } L_{a,T,1} < \ell \leq L_{a,T,2}, \\ 0, & \text{if } \ell > L_{a,T,2}, \end{cases} \quad (\text{G.40})$$

where $L_{a,T,1} = \frac{5c_1^2 |\mathcal{V}| \log T}{(\tilde{\Delta}_{\min}^a)^2}$, $L_{a,T,2} = \frac{10c_1^2 |\mathcal{V}| (m|\mathcal{V}| + |\mathcal{V}|) \log T}{(\tilde{\Delta}_{\min}^a)^2}$, and $\tilde{\Delta}_{\min}^a = \min_{k: p_a^{D,k} > 0, \Delta_k > 0} \Delta_k / p_a^{D,k}$. Then, the reason why Eq. (G.39) holds is proved as follows.

Proof.

1. When $T_{a,t-1} > L_{a,T,2} = \frac{10c_1^2 |\mathcal{V}| (m|\mathcal{V}| + |\mathcal{V}|) \log T}{(\tilde{\Delta}_{\min}^a)^2}$,

we have (G.39, a) $\leq \frac{10c_1^2 |\mathcal{V}| \log T}{T_{a,t-1} \cdot \tilde{\Delta}_{k_t}} - \frac{\tilde{\Delta}_{k_t}}{m|\mathcal{V}| + |\mathcal{V}|} < \frac{(\tilde{\Delta}_{\min}^a)^2}{(m|\mathcal{V}| + |\mathcal{V}|) \tilde{\Delta}_{k_t}} - \frac{\tilde{\Delta}_{k_t}}{m|\mathcal{V}| + |\mathcal{V}|} \leq 0 = \kappa_{a,T}(T_{a,t-1})$.

2. When $L_{a,T,1} < T_{a,t-1} \leq L_{a,T,2}$,

we have (G.39, a) $\leq \frac{10c_1^2 |\mathcal{V}| \log T}{T_{a,t-1} \cdot \tilde{\Delta}_{k_t}} - \frac{\tilde{\Delta}_{k_t}}{m|\mathcal{V}| + |\mathcal{V}|} < \frac{10c_1^2 |\mathcal{V}| \log T}{T_{a,t-1} \cdot \tilde{\Delta}_{k_t}} \leq \frac{10c_1^2 |\mathcal{V}| \log T}{T_{a,t-1} \cdot \tilde{\Delta}_{\min}^a} = \kappa_{a,T}(T_{a,t-1})$.

3. When $T_{a,t-1} \leq L_{a,T,1}$,

we further consider two different cases $T_{a,t-1} \leq \frac{5c_1^2 |\mathcal{V}| \log T}{(\tilde{\Delta}_{k_t})^2}$ or $\frac{5c_1^2 |\mathcal{V}| \log T}{(\tilde{\Delta}_{k_t})^2} < T_{a,t-1} \leq L_{a,T,1} = \frac{5c_1^2 |\mathcal{V}| \log T}{(\tilde{\Delta}_{\min}^a)^2}$.

- For the former case, if there exists $a \in \tau_t$ so that $T_{a,t-1} \leq \frac{5c_1^2 |\mathcal{V}| \log T}{(\tilde{\Delta}_{k_t})^2}$, then we know $\sum_{q \in \tilde{\mathcal{S}}_t} \kappa_{q,T}(T_{t-1,q}) \geq \kappa_{a,T}(T_{a,t-1}) =$

$2\sqrt{\frac{5c_1^2 |\mathcal{V}| \log T}{T_{a,t-1}}} \geq 2\sqrt{(\tilde{\Delta}_{k_t})^2} \geq \Delta_{k_t}$, which makes Eq. (G.39) valid no matter what.

- For the latter case, when $\frac{5c_1^2 |\mathcal{V}| \log T}{(\tilde{\Delta}_{k_t})^2} < T_{a,t-1}$, we know that (G.39, a) $\leq \frac{10c_1^2 |\mathcal{V}| \log T}{\tilde{\Delta}_{k_t}} \frac{1}{T_{a,t-1}} = 2\sqrt{\frac{5c_1^2 |\mathcal{V}| \log T}{(\tilde{\Delta}_{k_t})^2}} \frac{1}{T_{a,t-1}} \sqrt{\frac{5c_1^2 |\mathcal{V}| \log T}{T_{a,t-1}}} \leq$

$2\sqrt{\frac{\tilde{\Delta}_{k_t}}{\tilde{\Delta}_{k_t}}} \sqrt{\frac{5c_1^2 |\mathcal{V}| \log T}{T_{a,t-1}}} = 2\sqrt{\frac{5c_1^2 |\mathcal{V}| \log T}{T_{a,t-1}}} = \kappa_{a,T}(T_{a,t-1})$.

4. When $\ell = 0$,

We have (G.39), $a \leq \frac{10c_1^2|\mathcal{V}|}{\Delta_{k_t}} \cdot \frac{1}{28} - \frac{\tilde{\Delta}_{k_t}}{m|\mathcal{V}| + |\mathcal{V}|} \leq \frac{1.25c_1^2|\mathcal{V}|}{\Delta_{k_t}} \leq \frac{1.25c_1^2|\mathcal{V}|}{\tilde{\Delta}_{\min}^a} = \kappa_{a,T}(T_{a,t-1})$.

Combining all the above cases, we have $\Delta_{k_t} \leq \mathbb{E}[\sum_{a \in \tau_t} \kappa_{a,T}(T_{a,t-1})]$. $\square$

Based on Eq. (G.39), we now have the following inequalities on $\text{Reg}_{\mu,\sigma}(T, E_{t,1})$:

$$\text{Reg}_{\mu,\sigma}(T, E_{t,1}) = \mathbb{E} \left[\sum_{t=1}^T \Delta_{k_t} \right] \quad (\text{G.41})$$

$$\stackrel{(a)}{\leq} \mathbb{E} \left[\sum_{t \in [T]} \mathbb{E}_t \left[\sum_{a \in \tau_t} \kappa_{a,T}(T_{a,t-1}) \right] \right] \quad (\text{G.42})$$

$$\stackrel{(b)}{=} \mathbb{E} \left[\sum_{t \in [T]} \sum_{a \in \tau_t} \kappa_{a,T}(T_{a,t-1}) \right] \quad (\text{G.43})$$

$$\stackrel{(c)}{=} \mathbb{E} \left[\sum_{a \in \mathcal{A} \cup \mathcal{V}} \sum_{s=0}^{T_{a,t-1}} \kappa_{a,T}(s) \right] \quad (\text{G.44})$$

$$\leq \sum_{a \in \mathcal{A} \cup \mathcal{V}} \frac{1.25c_1^2|\mathcal{V}|}{\kappa_{a,T}} + \sum_{a \in \mathcal{A} \cup \mathcal{V}} \sum_{s=1}^{L_{a,T,1}} 2\sqrt{\frac{5c_1^2|\mathcal{V}| \log T}{s}} \\ + \sum_{a \in \mathcal{A} \cup \mathcal{V}} \sum_{s=L_{a,T,1}+1}^{L_{a,T,2}} \frac{10c_1^2|\mathcal{V}| \log T}{\kappa_{a,T}} \frac{1}{s} \quad (\text{G.45})$$

$$\leq \sum_{a \in \mathcal{A} \cup \mathcal{V}} \frac{1.25c_1^2|\mathcal{V}|}{\tilde{\Delta}_{\min}^a} + \sum_{a \in \mathcal{A} \cup \mathcal{V}} \int_{s=0}^{L_{a,T,1}} 2\sqrt{\frac{5c_1^2|\mathcal{V}| \log T}{s}} ds \\ + \sum_{a \in \mathcal{A} \cup \mathcal{V}} \int_{s=L_{a,T,1}}^{L_{a,T,2}} \frac{10c_1^2|\mathcal{V}| \log T}{\tilde{\Delta}_{\min}^a s} ds \quad (\text{G.46})$$

$$\leq \sum_{a \in \mathcal{A} \cup \mathcal{V}} \frac{1.25c_1^2|\mathcal{V}|}{\tilde{\Delta}_{\min}^a} + \sum_{a \in \mathcal{A} \cup \mathcal{V}} \frac{10c_1^2|\mathcal{V}| \log T}{\tilde{\Delta}_{\min}^a} (1 + \log(m|\mathcal{V}| + |\mathcal{V}|)) \\ + \sum_{a \in \mathcal{A} \cup \mathcal{V}} \frac{20c_1^2|\mathcal{V}| \log T}{\tilde{\Delta}_{\min}^a}, \quad (\text{G.47})$$

where (a) follows from Eq. (G.39), (b) follows from the tower rule, (c) follows from that $T_{a,t-1}$ is increased by 1 if and only if $a \in \tau_t$.

G.3.2. Bounding the regret on the second term

Let $c_2 = 28$. Given filtration $\mathcal{F}_{t-1}$ and event $E_{t,2}$, we have

$$\Delta_{k_t} \stackrel{(a)}{\leq} \sum_{a \in \tilde{S}_t} 2c_2 p_a^{D,k_t} \min \left\{ 1/28, \frac{\log T}{T_{a,t-1}} \right\} \\ \stackrel{(b)}{\leq} -\Delta_{k_t} + 2 \sum_{a \in \tilde{S}_t} 2c_2 p_a^{D,k_t} \min \left\{ 1/28, \frac{\log T}{T_{a,t-1}} \right\} \\ = -\frac{\sum_{a \in \mathcal{A} \cup \mathcal{V}} p_a^{D,k_t} \Delta_{k_t}}{\sum_{a \in \mathcal{A} \cup \mathcal{V}} p_a^{D,k_t}} + 2 \sum_{i \in [m]} 2c_2 p_a^{D,k_t} \min \left\{ 1/28, \frac{\log T}{T_{a,t-1}} \right\} \quad (\text{G.48})$$

$$\stackrel{(c)}{\leq} \sum_{a \in \mathcal{A} \cup \mathcal{V}} p_a^{D,k_t} \left(-\frac{\Delta_{k_t}}{m|\mathcal{V}| + |\mathcal{V}|} + 4c_2 \min \left\{ 1/28, \frac{\log T}{T_{a,t-1}} \right\} \right), \quad (\text{G.49})$$

where (a) follows from event $E_{t,2}$, (b) is by the reverse amortization trick [17] that multiplies two to both sides of (a) and rearranges the terms, and (c) follows from $p_a^{D,k_t} \leq 1$.

It follows that

$$\Delta_{k_t} = \mathbb{E}_t[\Delta_{k_t}] \stackrel{(a)}{\leq} \mathbb{E}_t \left[\sum_{a \in \mathcal{A} \cup \mathcal{V}} p_a^{D,k_t} \left(-\frac{\Delta_{k_t}}{m|\mathcal{V}| + |\mathcal{V}|} + 4c_2 \min \left\{ 1/28, \frac{\log T}{T_{a,t-1}} \right\} \right) \right] \\ \stackrel{(b)}{=} \mathbb{E}_t \left[\sum_{a \in \tau_t} \left(-\frac{\Delta_{k_t}}{m|\mathcal{V}| + |\mathcal{V}|} + 4c_2 \min \left\{ 1/28, \frac{\log T}{T_{a,t-1}} \right\} \right) \right]$$

$$\stackrel{(c)}{\leq} \mathbb{E}_t \left[\sum_{a \in \tau_t} \kappa_{a,T}(T_{a,t-1}) \right], \quad (\text{G.50})$$

Here, the regret allocation function is defined as

$$\kappa_{a,T}(\ell) = \begin{cases} \Delta_{\max}^a, & \text{if } 0 \leq \ell \leq L_{a,T,1} \\ \frac{4c_2 \log T}{\ell}, & \text{if } L_{a,1} < \ell \leq L_{a,2} \\ 0, & \text{if } \ell > L_{a,T,2} + 1, \end{cases} \quad (\text{G.51})$$

where $L_{a,T,1} = \frac{4c_2 \log T}{\Delta_{\max}^a}$, $L_{a,T,2} = \frac{4c_2(m|\mathcal{V}|+|\mathcal{V}|)\log T}{\Delta_{\min}^a}$. And (a) follows from Eq. (G.49), (b) from the triggering probability equivalence technique [40] that links the triggering probabilities with the random triggering event under expectation, and (c) is proven similarly to Section G.3.1 as follows.

Proof.

1. When $T_{a,t-1} > L_{a,T,2} = \frac{4c_2(m|\mathcal{V}|+|\mathcal{V}|)\log T}{\Delta_{\min}^a}$,
we have (G.50, i) $\leq 4c_2 \frac{\log T}{T_{a,t-1}} - \frac{\Delta_{k_t}}{m|\mathcal{V}|+|\mathcal{V}|} < \frac{\Delta_{\min}^a}{m|\mathcal{V}|+|\mathcal{V}|} - \frac{\Delta_{k_t}}{m|\mathcal{V}|+|\mathcal{V}|} \leq 0 = \kappa_{a,T}(T_{a,t-1})$.
2. When $T_{a,t-1} \leq L_{a,T,2}$,
we have (G.50, i) $\leq 4c_2 \frac{\log T}{T_{a,t-1}} - \frac{\Delta_{k_t}}{m|\mathcal{V}|+|\mathcal{V}|} < \frac{4c_2 \log T}{T_{a,t-1}} = \kappa_{a,T}(T_{a,t-1})$.
3. When $T_{a,t-1} \leq L_{a,T,1}$,
if there exists $a \in \tilde{S}_t$ so that $T_{a,t-1} \leq L_{a,T,1}$, then we know $\sum_{q \in \tilde{S}_t} \kappa_{q,T}(T_{t-1,q}) \geq \kappa_{a,T}(T_{a,t-1}) = \Delta_{\max}^a \geq \Delta_{k_t}$, which makes Eq. (G.50) holds no matter what. This means we do not need to consider this case for good.
Combining all above cases, we have $\Delta_{k_t} \leq \mathbb{E}_t[\sum_{a \in \tau_t} \kappa_{a,T}(T_{a,t-1})]$. $\square$

$$\text{Reg}_{\mu,\sigma}(T, E_{t,2}) = \mathbb{E} \left[\sum_{t=1}^T \Delta_{k_t} \right] \quad (\text{G.52})$$

$$\stackrel{(a)}{\leq} \mathbb{E} \left[\sum_{t \in [T]} \mathbb{E}_t \left[\sum_{a \in \tau_t} \kappa_{a,T}(T_{a,t-1}) \right] \right] \quad (\text{G.53})$$

$$\stackrel{(b)}{=} \mathbb{E} \left[\sum_{t \in [T]} \sum_{a \in \tau_t} \kappa_{a,T}(T_{a,t-1}) \right] \quad (\text{G.54})$$

$$\stackrel{(c)}{=} \mathbb{E} \left[\sum_{a \in \mathcal{A} \cup \mathcal{V}} \sum_{s=0}^{T_{a,t-1}} \kappa_{a,T}(s) \right] \quad (\text{G.55})$$

$$\leq |\mathcal{A} \cup \mathcal{V}| \Delta_{\max} + \sum_{a \in \mathcal{A} \cup \mathcal{V}} \sum_{\ell=1}^{L_{a,T,1}} \Delta_{\max}^a + \sum_{a \in \mathcal{A} \cup \mathcal{V}} \sum_{\ell=L_{a,T,1}+1}^{L_{a,T,2}} \frac{4c_2 \log T}{\ell} \quad (\text{G.56})$$

$$\leq |\mathcal{A} \cup \mathcal{V}| \Delta_{\max} + \sum_{a \in \mathcal{A} \cup \mathcal{V}} 4c_2 \log T + \sum_{a \in \mathcal{A} \cup \mathcal{V}} 4c_2 \log \left(\frac{(m|\mathcal{V}|+|\mathcal{V}|)\Delta_{\max}^a}{\Delta_{\min}^a} \right) \log T \quad (\text{G.57})$$

$$= |\mathcal{A} \cup \mathcal{V}| \Delta_{\max} + \sum_{a \in \mathcal{A} \cup \mathcal{V}} 4c_2 \left(1 + \log \left(\frac{(m|\mathcal{V}|+|\mathcal{V}|)\Delta_{\max}^a}{\Delta_{\min}^a} \right) \right) \log T \quad (\text{G.58})$$

$$\leq |\mathcal{A} \cup \mathcal{V}| \Delta_{\max} + \sum_{a \in \mathcal{A} \cup \mathcal{V}} 4c_2 \left(1 + \log \left(\frac{(m|\mathcal{V}|+|\mathcal{V}|)\Delta_{\max}^a}{\Delta_{\min}^a} \right) \right) \log T, \quad (\text{G.59})$$

where (a) follows from Eq. (G.50), (b) follows from the tower rule, (c) follows from that $T_{a,t-1}$ is increased by 1 if and only if $a \in \tau_t$.

In the case of fixed node weights, the problem degenerates into one where only edge weights need to be learned, as the node weights remain constant and do not require learning.

G.3.3. Distribution-independent bound

We fix a gap Δ , and consider two events for Δ_{k_t} : $\{\Delta_{k_t} \leq \Delta\}$ and $\{\Delta_{k_t} > \Delta\}$. In the former case, the regret is trivially bounded by $\text{Reg}(T, \{\Delta_{k_t} \leq \Delta\}) \leq T\Delta$. For the latter case, under $\{\Delta_{k_t} > \Delta\}$, we can easily replace all instances of $\Delta_{\min}^a$ with Δ and derive the bound $\text{Reg}(T, \{\Delta_{k_t} > \Delta\}) \leq O\left(\frac{|\mathcal{A} \cup \mathcal{V}| \log(m|\mathcal{V}|+|\mathcal{V}|)\log T}{\Delta} + |\mathcal{A} \cup \mathcal{V}| \log \left(\frac{m|\mathcal{V}|+|\mathcal{V}|}{\Delta} \right) \log T\right)$. Similar to Section F.3, we can also establish a

distribution-independent bound:

$$O\left(\sqrt{|\mathcal{A} \cup \mathcal{V}| |\mathcal{V}| T \log(m|\mathcal{V}| + |\mathcal{V}|) \log T} + |\mathcal{A} \cup \mathcal{V}| \log T \log((m|\mathcal{V}| + |\mathcal{V}|)T)\right).$$

Appendix H. Online learning for non-overlapping MuLaNE

H.1. Online learning algorithm for non-overlapping MuLaNE

For non-overlapping MuLaNE, we estimate layer-level marginal gains as our parameters, rather than probabilities $P_{i,u}(b)$ and node weights σ . By doing so, we can largely simplify our analysis since we no longer need to handle the unknown weights and the extra constraint of $P_{i,u}(b)$ as in the overlapping setting. Concretely, we maintain a set of base arms $\mathcal{A} = \{(i, b) | i \in [m], b \in [c_i]\}$, and let $|\mathcal{A}| = \sum_{i \in [m]} c_i$ be the total number of these arms. For each base arm $(i, b) \in \mathcal{A}$, denote $\mu_{i,b}$ as the truth value of each base arm, i.e., $\mu_{i,b} = \sum_{u \in \mathcal{V}} \sigma_u (P_{i,u}(b) - P_{i,u}(b-1))$. Based on the definition of base arms, we can see that we do not need to explicitly estimate the node weights σ or probabilities $P_{i,u}(b)$, so we will use $r_\mu(k)$ to denote the reward function $r_{G,\alpha,\sigma}(k)$.

We present our algorithm in Algorithm 8, which has three differences compared with the Algorithm 3 for overlapping MuLaNE. First, the expected reward of action k is $r_\mu(k) = \sum_{i=1}^m \sum_{b=1}^{k_i} \mu_{i,b}$, which is a linear reward function and satisfies two new Monotonicity and 1-Norm Bounded Smoothness conditions (see Condition 1 and 2). Second, we use the Dynamic Programming method (Algorithm 7) as our Oracle, which does not require the parameters $\mu_{i,b}$ to be increasing. Moreover, the offline oracle always outputs *optimal* solutions. Third, we use a new random variable $Y_{i,b}$, which indicates the gain of weights which depends on whether a new node is visited by random walker from the i -th layer exactly at the b -th step, i.e., $Y_{i,b} = \sum_{u \in \bigcup_{j=1}^b \{X_{i,j}\}} \sigma_u - \sum_{u \in \bigcup_{j=1}^{b-1} \{X_{i,j}\}} \sigma_u$ for base arms (i, b) such that $b \leq k_i$.

Condition 1 (Monotonicity). The reward $r_\mu(k)$ is monotonically increasing, i.e., for any budget allocation k and any two vectors $\mu = (\mu_{i,b})_{(i,b) \in \mathcal{A}}$, $\mu' = (\mu'_{i,b})_{(i,b) \in \mathcal{A}}$, we have $r_\mu(k) \leq r_{\mu'}(k)$, if $\mu_{i,b} \leq \mu'_{i,b}$ for $(i, b) \in \mathcal{A}$.

Condition 2 (1-Norm Bounded Smoothness). The reward function $r_\mu(k)$ satisfies the 1-norm bounded smoothness, i.e., for any budget allocation k and any two vectors $\mu = (\mu_{i,b})_{(i,b) \in \mathcal{A}}$, $\mu' = (\mu'_{i,b})_{(i,b) \in \mathcal{A}}$, we have $|r_\mu(k) - r_{\mu'}(k)| \leq \sum_{b \leq k_i} |\mu_{i,b} - \mu'_{i,b}|$.

We define the reward gap $\Delta_k = r_\mu(k^*) - r_\mu(k)$ for all feasible action k satisfying $\sum_{i=1}^m k_i = B$, $0 \leq k_i \leq c_i$. For each base arm (i, b) , we define $\Delta_{\min}^{i,b} = \min_{\Delta_k > 0, k_i \geq b} \Delta_k$ and $\Delta_{\max}^{i,b} = \max_{\Delta_k > 0, k_i \geq b} \Delta_k$. As a convention, if there is no action with $k_i \geq b$ such that $\Delta_k > 0$, we define $\Delta_{\min}^{i,b} = \infty$ and $\Delta_{\max}^{i,b} = 0$. Also, we define $\Delta_{\min} = \min_{(i,b) \in \mathcal{A}} \Delta_{\min}^{i,b}$ and $\Delta_{\max} = \max_{(i,b) \in \mathcal{A}} \Delta_{\max}^{i,b}$.

Theorem 9. CUCB-MG has the following distribution-dependent regret bound,

$$\text{Reg}_\mu(T) \leq \sum_{(i,b) \in \mathcal{A}} \frac{48B \ln T}{\Delta_{\min}^{i,b}} + 2|\mathcal{A}| + \frac{\pi^2}{3} |\mathcal{A}| \Delta_{\max}.$$

Algorithm 8 CUCB-MG: combinatorial upper confidence bound (CUCB) algorithm for non-overlapping MuLaNE, using Marginal Gains as the base arms.

Input: Budget B , number of layers m , offline oracle DP.

- 1: For each $(i, b) \in \mathcal{A}$, $T_{i,b} \leftarrow 0$, $\hat{\mu}_{i,b} \leftarrow 1$.
- 2: **for** $t = 1, 2, 3, \dots, T$ **do**
- 3: **for** $(i, b) \in \mathcal{A}$ **do**
- 4: $\rho_{i,b} \leftarrow \sqrt{3 \ln t / (2T_{i,b})}$.
- 5: $\hat{\mu}_{i,b} = \min\{\hat{\mu}_{i,b} + \rho_{i,b}, 1\}$.
- 6: **end for**
- 7: $k \leftarrow \text{DP}((\hat{\mu}_{i,b})_{i,b}, B, c)$
- 8: Apply budget allocation k , which gives trajectories $X := (X_{i,1}, \dots, X_{i,k_i})_{i \in [m]}$ as feedbacks.
- 9: For $(i, b) \in \tau := \{(i, b) \in \mathcal{A} | b \leq k_i\}$,
- $Y_{i,b} = \sum_{u \in \bigcup_{j=1}^b \{X_{i,j}\}} \sigma_u - \sum_{u \in \bigcup_{j=1}^{b-1} \{X_{i,j}\}} \sigma_u$.
- 10: For $(i, b) \in \tau$, update $T_{i,b}$ and $\hat{\mu}_{i,b}$:
- 11: $T_{i,b} = T_{i,b} + 1$, $\hat{\mu}_{i,b} = \hat{\mu}_{i,b} + (Y_{i,b} - \hat{\mu}_{i,b}) / T_{i,b}$.
- 12: **end for**

H.2. Regret analysis for non-overlapping MuLaNE

In this section, we give the regret analysis for CUCB-MG, which applies the standard CUCB algorithm [17] to this setting.

Let $\mathcal{A}$ denote the set containing all base arms, i.e., $\mathcal{A} = \{(i, b) | i \in [m], b \in [c_i]\}$. We add subscript t to denote the value of a variable at the end of t . For example, $T_{i,b,t}$ denotes the total times of arm $(i, b) \in \mathcal{A}$ is played at the end of round t . Let us first introduce a definition of an unlikely event that $\hat{\mu}_{i,b,t-1}$ is not accurate as expected.

Definition 5. We say that the sampling is nice at the beginning of round t , if for every arm $(i, b) \in \mathcal{A}$, $|\hat{\mu}_{i,b,t-1} - \mu_{i,b}| < \rho_{i,b,t}$, where $\rho_{i,b,t} = \sqrt{\frac{3 \ln t}{2T_{i,b,t-1}}}$. Let $\mathcal{N}_t$ be such event.

Lemma 18. For each round $t \geq 1$, $Pr\{\neg \mathcal{N}_t\} \leq 2|\mathcal{A}|t^{-2}$

Proof. For each round $t \geq 1$, we have

$$\begin{aligned}
& Pr\{\neg \mathcal{N}_t\} \\
&= Pr\{\exists (i, b) \in \mathcal{A}, |\hat{\mu}_{i,b,t-1} - \mu_{i,b}| \geq \sqrt{\frac{3 \ln t}{2T_{i,b,t-1}}}\} \\
&\leq \sum_{(i,b) \in \mathcal{A}} Pr\{|\hat{\mu}_{i,b,t-1} - \mu_{i,b}| \geq \sqrt{\frac{3 \ln t}{2T_{i,b,t-1}}}\} \\
&= \sum_{(i,b) \in \mathcal{A}} \sum_{k=1}^{t-1} Pr\{T_{i,b,t-1} = k \\
&\quad \wedge |\hat{\mu}_{i,b,t-1} - \mu_{i,b}| \geq \sqrt{\frac{3 \ln t}{2T_{i,b,t-1}}}\} \\
&\leq \sum_{(i,b) \in \mathcal{A}} \sum_{k=1}^{(t-1)} \frac{2}{k^3} \leq 2|\mathcal{A}|t^{-2}.
\end{aligned}$$

When $T_{i,b,t-1} = k$, $\hat{\mu}_{i,b,t-1}$ is the average of k i.i.d. random variables $Y_{i,b}^{[1]}, \dots, Y_{i,b}^{[k]}$, where $Y_{i,b}^{[j]}$ is the Bernoulli random variable of arm (i, b) when it is played for the j -th time during the execution. That is, $\hat{\mu}_{i,b,t-1} = \sum_{j=1}^k Y_{i,b}^{[j]} / k$. Note that the independence of $Y_{i,b}^{[1]}, \dots, Y_{i,b}^{[k]}$ comes from the fact we select a new starting position following the starting distribution α_t at the beginning of each round t . The last inequality uses the Hoeffding Inequality [37]. $\square$

Then we can use monotonicity and 1-norm bounded smoothness properties (Conditions 1 and 2) to bound the reward gap $\Delta_{k_t} = r_{\mu}(k^*) - r_{\mu}(k_t)$ between the optimal action k^* and the action $k_t = (k_{t,1}, \dots, k_{t,m})$ selected by our algorithm A . To achieve this, we introduce a positive real number $M_{i,b}$ for each arm (i, b) and define $M_{k_t} = \max_{(i,b) \in S_t} M_{i,b}$, where $S_t = \{(i, b) \in \mathcal{A} | b = k_{t,i}\}$ and $|S_t| = B$. Define,

$$\kappa_T(M, s) = \begin{cases} 2, & \text{if } s = 0, \\ 2\sqrt{\frac{6 \ln T}{s}}, & \text{if } 1 \leq s \leq \ell_T(M), \\ 0, & \text{if } s \geq \ell_T(M) + 1, \end{cases}$$

where

$$\ell_T(M) = \lfloor \frac{24B^2 \ln T}{M^2} \rfloor.$$

Lemma 19. If $\{\Delta_{k_t} \geq M_{k_t}\}$ and $\mathcal{N}_t$ holds, we have

$$\Delta_{k_t} \leq \sum_{(i,b) \in S_t} \kappa_T(M_{i,b}, T_{i,b,t-1}). \quad (\text{H.1})$$

Proof.

First, we can observe $\bar{\mu}_{i,b,t} \geq \mu_{i,b}$, for $(i, b) \in \mathcal{A}$, when $\mathcal{N}_t$ holds. This comes from the fact $\bar{\mu}_{i,b,t} = \min\{\hat{\mu}_{i,b,t-1} + \rho_{i,b,t}, 1\}$ and $|\hat{\mu}_{i,b,t-1} - \mu_{i,b}| < \rho_{i,b,t}$.

Notice that the right hand of Eq. (H.1) is non-negative, it is trivially satisfied when $\Delta_{k_t} = 0$. We only need to consider $\Delta_{k_t} > 0$. By $\mathcal{N}_t$ and Condition 1, We have the following inequalities,

$$r_{\bar{\mu}_t}(k_t) \geq r_{\bar{\mu}_t}(k^*) \geq r_{\mu}(k^*) = \Delta_{k_t} + r_{\mu}(k_t).$$

The first inequality is due to the oracle outputs the optimal solution given parameters $\bar{\mu}_t$, the second inequality is due to Condition 1 (monotonicity).

Then, by Condition 2, we have

$$\Delta_{k_t} \leq r_{\bar{\mu}_t}(k_t) - r_{\mu}(k_t) \leq \sum_{(i,b) \in S_t} |\bar{\mu}_{i,b,t} - \mu_{i,b}|.$$

Next, we can bound Δ_{k_t} by bounding $\bar{\mu}_{i,b,t} - \mu_{i,b}$ by applying a transformation. As $\{\Delta_{k_t} \geq M_{k_t}\}$ holds by assumption, so $\sum_{(i,b) \in S_t} |\bar{\mu}_{i,b,t} - \mu_{i,b}| \geq \Delta_{k_t} \geq M_{k_t}$. We have

$$\Delta_{k_t} \leq \sum_{(i,b) \in S_t} |\bar{\mu}_{i,b,t} - \mu_{i,b}|$$

$$\begin{aligned}
&\leq -M_{k_t} + 2 \sum_{(i,b) \in S_t} |\bar{\mu}_{i,b,t} - \mu_{i,b}| \\
&= 2 \sum_{(i,b) \in S_t} (|\bar{\mu}_{i,b,t} - \mu_{i,b}| - \frac{M_{k_t}}{2B}) \\
&\leq 2 \sum_{(i,b) \in S_t} (|\bar{\mu}_{i,b,t} - \mu_{i,b}| - \frac{M_{i,b}}{2B}).
\end{aligned} \tag{H.2}$$

By $\mathcal{N}_t$, we have:

$$\begin{aligned}
|\bar{\mu}_{i,b,t} - \mu_{i,b}| - \frac{M_{i,b}}{2B} &\leq \min\{2\rho_{i,b,t}, 1\} - \frac{M_{i,b}}{2B} \\
&\leq \min\left\{2\sqrt{\frac{3 \ln T}{2T_{i,b,t-1}}}, 1\right\} - \frac{M_{i,b}}{2B}.
\end{aligned}$$

Now consider the following cases:

Case 1. If $T_{i,b,t-1} \leq \ell_T(M_{i,b})$, we have $\bar{\mu}_{i,b,t} - \mu_{i,b} - \frac{M_{i,b}}{2B} \leq \min\left\{2\sqrt{\frac{3 \ln T}{2T_{i,b,t-1}}}, 1\right\} \leq \frac{1}{2}\kappa_T(M_{i,b}, T_{i,b,t-1})$.

Case 2. If $T_{i,b,t-1} \geq \ell_T(M_{i,b}) + 1$, then $2\sqrt{\frac{3 \ln T}{2T_{i,b,t-1}}} \leq \frac{M_{i,b}}{2B}$, so $\bar{\mu}_{i,b,t} - \mu_{i,b} - \frac{M_{i,b}}{2B} \leq 0 = \frac{1}{2}\kappa_T(M_{i,b}, T_{i,b,t-1})$.

Combining these two cases and Eq. (H.2), we have the following inequality, which concludes the lemma.

$$\Delta_{k_t} \leq \sum_{(i,b) \in S_t} \kappa_T(M_{i,b}, T_{i,b,t-1}). \quad \square$$

Proof. By above lemmas, we have

$$\begin{aligned}
&\sum_{t=1}^T \mathbb{I}(\{\Delta_{k_t} \geq M_{k_t}\} \wedge \mathcal{N}_t) \Delta_{k_t} \\
&\leq \sum_{t=1}^T \sum_{(i,b): (i,b) \in S_t} \kappa_T(M_{i,b}, T_{i,b,t-1}) \\
&= \sum_{(i,b) \in \mathcal{A}} \sum_{t' \in \{t | t \in [T], k_{t,i} \geq b\}} \kappa_T(M_{i,b}, T_{i,b,t'-1}) \\
&\leq \sum_{(i,b) \in \mathcal{A}} \sum_{s=0}^{T_{i,b,T}} \kappa_T(M_{i,b}, s) \\
&\leq 2|\mathcal{A}| + \sum_{(i,b) \in \mathcal{A}} \sum_{s=1}^{\ell_T(M_{i,b})} 2\sqrt{6 \frac{\ln T}{s}} \\
&\leq 2|\mathcal{A}| + \sum_{(i,b) \in \mathcal{A}} \int_{s=0}^{\ell_T(M_{i,b})} 2\sqrt{6 \frac{\ln T}{s}} ds \\
&\leq 2|\mathcal{A}| + \sum_{(i,b) \in \mathcal{A}} 4\sqrt{6 \ln T \ell_T(M_{i,b})} \\
&\leq 2|\mathcal{A}| + \sum_{(i,b) \in \mathcal{A}} \frac{48B \ln T}{M_{i,b}}.
\end{aligned}$$

Let us set $M_{i,b} = \Delta_{\min}^{i,b}$ and thus the event $\{\Delta_{k_t} \geq M_{k_t}\}$ always holds. We have

$$\begin{aligned}
&\sum_{t=1}^T \mathbb{I}(\mathcal{N}_t \wedge \{\Delta_{k_t} \geq M_{k_t}\}) \Delta_{k_t} \\
&\leq 2|\mathcal{A}| + \sum_{(i,b) \in \mathcal{A}} \frac{48B \ln T}{\Delta_{\min}^{i,b}}.
\end{aligned}$$

By Lemma 18, $Pr\{\neg \mathcal{N}_t\} \leq 2|\mathcal{A}|t^{-2}$, then we have

$$\mathbb{E}\left[\sum_{t=1}^T \mathbb{I}(\neg \mathcal{N}_t) \Delta_{k_t}\right] \leq \sum_{t=1}^T 2|\mathcal{A}|t^{-2} \Delta_{\max} \leq \frac{\pi^2}{3} |\mathcal{A}| \Delta_{\max}.$$

By our choice of $M_{i,b}$,

$$\sum_{t=1}^T \mathbb{I}(\Delta_{k_t} < M_{k_t}) \Delta_{k_t} = 0.$$

By the definition of the regret,

$$\begin{aligned}
 \text{Reg}_\mu(T) &= \mathbb{E} \left[\sum_{t=1}^T \Delta_{k_t} \right] \\
 &\leq \mathbb{E} \left[\sum_{t=1}^T \left(\mathbb{I}(\neg \mathcal{N}_t) + \mathbb{I}(\Delta_{k_t} < M_{k_t}) \right. \right. \\
 &\quad \left. \left. + \mathbb{I}(\mathcal{N}_t \wedge \{\Delta_{k_t} \geq M_{k_t}\}) \right) \Delta_{k_t} \right] \\
 &\leq \sum_{(i,b) \in \mathcal{A}} \frac{48B \ln T}{\Delta_{i,b}^{\min}} + 2|\mathcal{A}| + \frac{\pi^2}{3} |\mathcal{A}| \Delta_{\max}.
 \end{aligned}$$

For distribution independent bound, we take $M_{i,b} = M = \sqrt{48B|\mathcal{A}| \ln T/T}$, then we have $\sum_{t=1}^T \mathbb{I}\{\Delta_{k_t} \leq M_{k_t}\} \leq TM$. Then

$$\begin{aligned}
 \text{Reg}_\mu(T) &\leq \sum_{(i,b) \in \mathcal{A}} \frac{48B \ln T}{M_{i,b}} + 2|\mathcal{A}| + \frac{\pi^2}{3} |\mathcal{A}| \Delta_{\max} + TM \\
 &\leq \frac{48B|\mathcal{A}| \ln T}{M} + 2|\mathcal{A}| + \frac{\pi^2}{3} |\mathcal{A}| \Delta_{\max} + TM \\
 &= 2\sqrt{48B|\mathcal{A}|T \ln T} + \frac{\pi^2}{3} |\mathcal{A}| \Delta_{\max} + 2|\mathcal{A}| \\
 &\leq 14\sqrt{B|\mathcal{A}|T \ln T} + \frac{\pi^2}{3} |\mathcal{A}| \Delta_{\max} + 2|\mathcal{A}|. \quad \square
 \end{aligned}$$

Appendix I. Supplemental experiments

I.1. Experimental results for stationary distributions

In this section, we show experimental results with explorers starting from stationary distributions in Fig. I.8. Specifically, each explorer W_i starts from the node u in layer L_i with probability proportional to the out-degree of u . For offline overlapping case, the total weights of unique nodes visited for BEG and MG are the same, and outperforms two baselines, which accords with our theoretical analysis. Similar to the Section 7, BEG and MG are empirically close to the optimal solution. For offline non-overlapping case, DP and MG give us the same *optimal* solution. For the online learning algorithms, note that we still use BEG rather than MG as the offline oracle, this is due to the UCB values do not have the diminishing return properties. The same applies for the non-overlapping case, where DP (instead of MG) is used as the offline oracle. The trend of the curves and the analysis are similar to that in the Section 7.

I.2. Online learning algorithms with different total budgets

As shown in Figs. I.9 and I.10, although there are some fluctuations, the trend of the curves and the analysis are almost the same as that when $B = 3000$, which shows our experimental results are consistent for different budgets.

I.3. Computational efficiency of online learning algorithms

For the computational efficiency of online learning algorithms, since the running time for algorithms with the same oracle is similar, we first present the running time for CUCB-MAX (or CUCB-MG) with different oracles in Table I.4a under $B = 3.0k$. CUCB-MAX with BEG is 50 times faster than BEGE, which is consistent with our theoretical analysis. Table I.4b demonstrates the total running time of CUCB-MAX and ECUCB-MAX algorithms with the same BEG oracle on the RocketFuel dataset under $T = 10000$, where the running time of ECUCB-MAX is slightly longer than that of CUCB-MAX. This is expected because the variance-based adjustments in ECUCB-MAX introduce extra computational complexity, although the difference in runtime remains relatively minor.

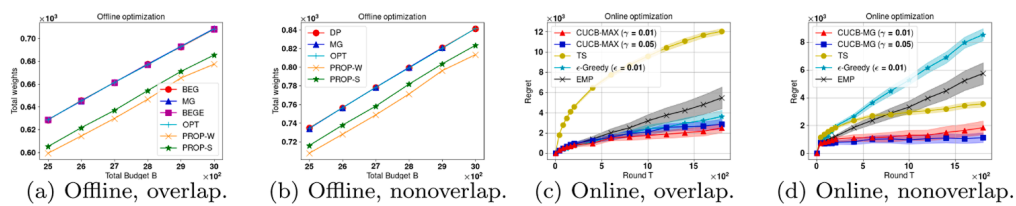


Fig. I.8. Explorers start with stationary distributions. Above: total weights of unique nodes visited given by different offline algorithms. Below: regret for different online algorithms when $B = 3000$.

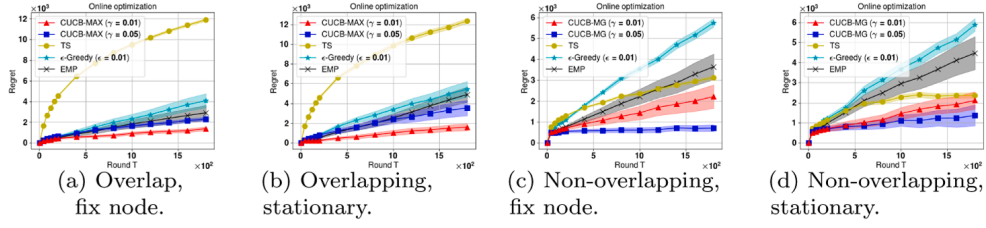
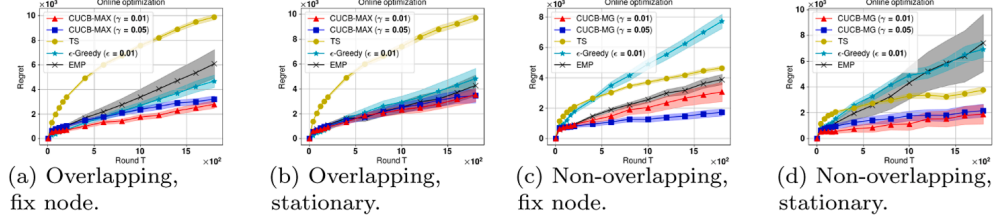
Fig. I.9. Regret for different online algorithms, when total budget $B = 2000$.Fig. I.10. Regret for different online algorithms, when total budget $B = 4000$.

Table I.4

Running time (seconds) for online algorithms on two different real-world datasets.

(a) With different offline oracles on the social network dataset.				
Scenario	BEG	BEGE	DP	OPT
Overlapping	1.22	60.03	NA	107.33
Non-overlapping	NA	NA	0.86	82.01
(b) With different parameters on the computer network dataset.				
Algorithm	$\gamma = 0.01$	$\gamma = 0.05$		
CUCB-MAX	97.6	96.9		
ECUCB-MAX	102.7	98.6		

I.4. Extended experimental settings and results

We extend the baseline experiments still on the FF-TW-YT dataset with its topology, by introducing more diverse weight assignments and randomized starting positions.

Settings. *Edge weights:* Instead of fixing all values to 1, we adopt two alternative schemes: (i) sample each edge weight independently from $\{0, 0.5, 1\}$ with equal probability, so that the network contains both weak, medium, and strong ties, reflecting heterogeneous interaction strengths among users; (ii) set the weight of edge $(u, v) \in \mathcal{E}$ to $\deg(u) + \deg(v)$, normalized by the maximum value in the network, where $\deg(\cdot)$ denotes the degree (number of incident edges) of a node. This scheme assigns larger weights to edges connecting highly connected users, modeling the intuition that interactions between central individuals are stronger or more influential. *Node weights:* Beyond the uniform random assignment $\{0, 0.5, 1\}$, we consider two alternatives: (i) *power-law weights*, where each node weight is drawn from a distribution $P(\sigma) \propto \sigma^{-\alpha}$ with $\alpha = 2.5$ and then linearly normalized to $[0, 1]$. This captures the skewness observed in real social systems, where a few nodes (e.g., influencers) are disproportionately important while the majority remain relatively unimportant; (ii) *centrality-correlated weights*, where $\sigma_u = \deg(u) / \max_{v \in \mathcal{V}} \deg(v)$, such that nodes with higher degree receive proportionally larger importance, mimicking the common assumption that well-connected users exert greater influence. For clarity, we denote the four settings as: (a) *Discrete edges & Power-law nodes*, (b) *Discrete edges & Degree-based nodes*, (c) *Degree-sum edges & Power-law nodes*, and (d) *Degree-sum edges & Degree-based nodes*. *Starting nodes:* Instead of always using the smallest node-id, the initial node in each trial is drawn uniformly at random from the vertex set of each layer. This avoids bias from a fixed entry point and better reflects realistic exploration where the walker may start from arbitrary users.

Baselines. We newly include CUCB [7] and BCUCB-T [17] as stronger baselines in our comparison. Notably, both [7, 17] propose the same CUCB algorithm, with [17] providing improved regret bounds via new proof techniques. BCUCB-T [28] is designed for settings with probabilistically triggered arms and achieves batch-size independence.

Results. We summarize the outcomes in Fig. I.11, where each subfigure reports regret at $T = 1500$ for the four algorithms under one weight configuration (panels (a)–(d)). Overall, ECUCB-MAX consistently achieves the lowest regret, where the gains stem from its property of precision-balanced bounded smoothness and explorer-relaxed design. CUCB-MAX performs reliably as the second-best method, while CUCB and BCUCB-T incur higher regret under heterogeneous conditions. Both rely on properties such as monotonically increasing UCB estimates, which are not guaranteed to hold in the budget-constrained setting, leading to slower adaptation, especially

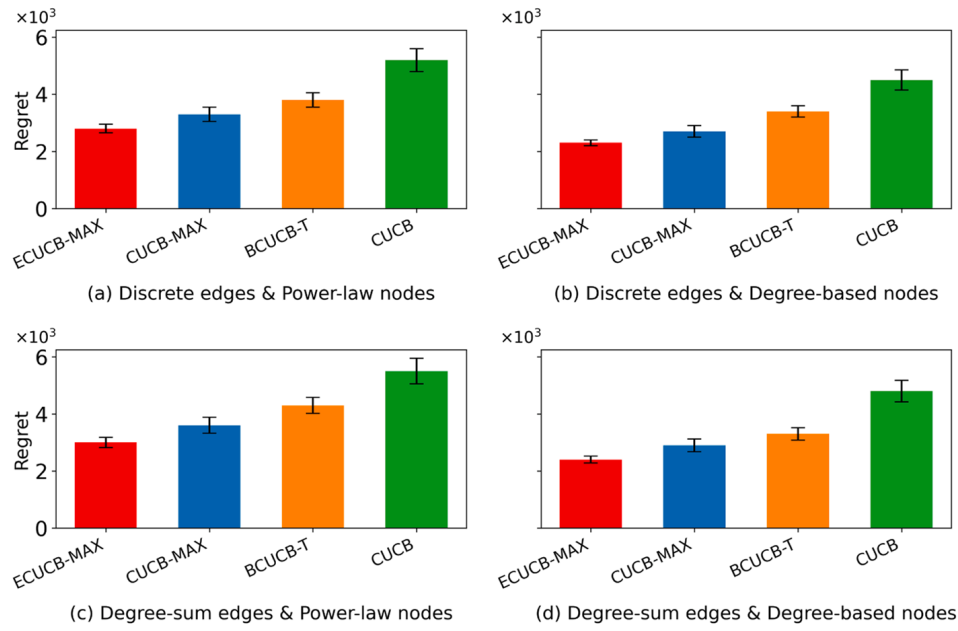


Fig. I.11. Regret@T=1500 of algorithms under extended experimental settings.

under heavy-tailed weight distributions. In the degree-correlated cases (panels (b) and (d)), absolute regret decreases for all methods since importance aligns with structure, but the ranking remains consistent. Finally, randomizing the starting node introduces only minor additional variance and does not affect the relative ordering of algorithms.

References

- [1] Q. Lv, P. Cao, E. Cohen, K. Li, S. Shenker, Search and replication in unstructured peer-to-peer networks, in: Proceedings of the 16th International Conference on Supercomputing, the 16th international conference on Supercomputing, 2002, pp. 84–95.
- [2] D.F. Gleich, Pagerank beyond the web 57 (2015) 321–363.
- [3] B. Wilder, N. Immerlica, E. Rice, M. Tambe, Maximizing influence in an unknown social network, in: Thirty-Second AAAI Conference on Artificial Intelligence, 2018.
- [4] R. Wang, Y. Li, S. Lin, W. Wu, H. Xie, Y. Xu, J.C. Lui, Common neighbors matter: fast random walk sampling with common neighbor awareness, IEEE Trans. Knowl. Data Eng. 35 (5) (2022) 4570–4584.
- [5] K. Lerman, L. Jones, Social browsing on flickr, arXiv preprint:cs/0612047.
- [6] X. Dai, Z. Wang, J. Xie, X. Liu, J.C. Lui, Conversational recommendation with online learning and clustering on misspecified users, IEEE Trans. Knowl. Data Eng. 36 (12) (2024) 7825–7838.
- [7] W. Chen, Y. Wang, Y. Yuan, Q. Wang, Combinatorial multi-armed bandit and its extension to probabilistically triggered arms, J. Mach. Learn. Res. 17 (1) (2016) 1746–1778.
- [8] Y. Li, Z. Wu, S. Lin, H. Xie, M. Lv, Y. Xu, J.C. Lui, Walking with perception: efficient random walk sampling via common neighbor awareness, in: IEEE 35th International Conference on Data Engineering (ICDE), IEEE, 2019, pp. 962–973.
- [9] D. Kempe, J. Kleinberg, E. Tardos, Maximizing the spread of influence through a social network, in: Proceedings of the Ninth ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, the ninth ACM SIGKDD international conference on Knowledge discovery and data mining, 2003, pp. 137–146.
- [10] W. Chen, Y. Wang, S. Yang, Efficient influence maximization in social networks, in: Proceedings of the 15th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, the 15th ACM SIGKDD international conference on Knowledge discovery and data mining, 2009, pp. 199–208.
- [11] C. Wang, W. Chen, Y. Wang, Scalable influence maximization for independent cascade model in large-scale social networks, Data Min. Knowl. Discov. 25 (3) (2012) 545–576.
- [12] N. Alon, I. Gamzu, M. Tennenholtz, Optimizing budget allocation among channels and influencers, in: Proceedings of the 21st International Conference on World Wide Web, the 21st international conference on World Wide Web, ACM, 2012, pp. 381–388.
- [13] T. Soma, N. Kakimura, K. Inaba, K.-I. Kawarabayashi, Optimal budget allocation: theoretical guarantee and efficient algorithm, in: International Conference on Machine Learning, 2014, pp. 351–359.
- [14] H. Robbins, Some aspects of the sequential design of experiments, Bull. Am. Math. Soc. 58 (5) (1952) 527–535.
- [15] S. Bubeck, N. Cesa-Bianchi, Regret analysis of stochastic and nonstochastic multi-armed bandit problems, Found. Trends Mach. Learn. 5 (1) (2012) 1–122.
- [16] Y. Gai, B. Krishnamachari, R. Jain, Combinatorial network optimization with unknown variables: multi-armed bandits with linear rewards and individual observations, IEEE/ACM Transact. Network. (TON) 20 (5) (2012) 1466–1478.
- [17] Q. Wang, W. Chen, Improving regret bounds for combinatorial semi-bandits with probabilistically triggered arms and its applications, in: Advances in Neural Information Processing Systems, 2017, pp. 1161–1171.
- [18] X. Chen, W. Huang, W. Chen, J.C. Lui, Community exploration: from offline optimization to online learning, in: Advances in Neural Information Processing Systems, 2018, pp. 5474–5483.
- [19] X. Liu, J. Zuo, X. Chen, W. Chen, J.C. Lui, Multi-layered network exploration via random walks: from offline optimization to online learning, in: International Conference on Machine Learning, PMLR, 2021, pp. 7057–7066.
- [20] X. Dai, X. Sun, J. Zuo, X. Liu, J.C. Lui, A unified online-offline framework for co-branding campaign recommendations, in: Proceedings of the 31st ACM SIGKDD Conference on Knowledge Discovery and Data Mining V, the 31st ACM SIGKDD Conference on Knowledge Discovery and Data Mining V, 2025, pp. 427–438.
- [21] M. Kapralov, I. Post, J. Vondrák, Online submodular welfare maximization: greedy is optimal, in: Proceedings of the Twenty-Fourth Annual ACM-SIAM Symposium on Discrete Algorithms, the twenty-fourth annual ACM-SIAM symposium on Discrete algorithms, 2013, pp. 1216–1225.

- [22] T. Soma, Y. Yoshida, A generalization of submodular cover via the diminishing return property on the integer lattice, in: *Advances in Neural Information Processing Systems*, 2015, pp. 847–855.
- [23] X. Dai, X. Liu, J. Zuo, H. Xie, C. Joe-Wong, J.C.S. Lui, Variance-aware bandit framework for dynamic probabilistic maximum coverage problem with triggered or self-reliant arms, *IEEE Transact. Network.* 33 (3) (2025) 982–993.
- [24] X. Liu, X. Dai, X. Wang, M. Hajiesmaili, J.C. Lui, Combinatorial logistic bandits, in: *abstracts of the 2025*, in: *ACM SIGMETRICS International Conference on Measurement and Modeling of Computer Systems*, 2025, pp. 112–114.
- [25] S. Khuller, A. Moss, J.S. Naor, The budgeted maximum coverage problem, *Inf. Process. Lett.* 70 (1) (1999) 39–45.
- [26] X. Liu, X. Dai, J. Zuo, S. Wang, C. Joe-Wong, J.C. Lui, W. Chen, Offline learning for combinatorial multi-armed bandits, in: *International Conference on Machine Learning*, PMLR, 2025, pp. 38251–38289.
- [27] N. Merlis, S. Mannor, Tight lower bounds for combinatorial multi-armed bandits, in: *Conference on Learning Theory*, PMLR, 2020, pp. 2830–2857.
- [28] X. Liu, J. Zuo, S. Wang, C. Joe-Wong, J. Lui, W. Chen, Batch-size independent regret bounds for combinatorial semi-bandits with probabilistically triggered arms or independent arms, *Adv. Neural Inf. Process. Syst.* 35 (2022) 14904–14916.
- [29] S. Wang, W. Chen, Thompson sampling for combinatorial semi-bandits, in: *International Conference on Machine Learning*, 2018, pp. 5114–5122.
- [30] M.E. Dickison, M. Magnani, L. Rossi, *Multilayer Social Networks*, Cambridge University Press, 2016.
- [31] N. Spring, R. Mahajan, D. Wetherall, Measuring isp topologies with rocketfuel, *ACM SIGCOMM Comput. Commun. Rev.* 32 (4) (2002) 133–145.
- [32] S. Dolev, Y. Elovici, R. Puzis, Routing betweenness centrality, *J. ACM (JACM)* 57 (4) (2010) 1–27.
- [33] M. Kijima, *Markov Processes for Stochastic Modeling*, 6, CRC Press, 1997.
- [34] R. Serfozo, *Basics of Applied Stochastic Processes*, Springer Science & Business Media, 2009.
- [35] M. Raginsky, I. Sason, et al., Concentration of measure inequalities in information theory, communications, and coding, *Foundat. Trends Commun. Inform. Theory*, 10, 2013.
- [36] A. Krause, J. Leskovec, C. Guestrin, J. Vanbriesen, C. Faloutsos, Efficient sensor placement optimization for securing large water distribution networks, *J. Water Resour. Plann. Manage.* 134 (6) (2008) 516–526.
- [37] D.P. Dubhashi, A. Panconesi, *Concentration of Measure for the Analysis of Randomized Algorithms*, Cambridge University Press, 2009.
- [38] J.-Y. Audibert, R. Munos, C. Szepesvári, Exploration-exploitation tradeoff using variance estimates in multi-armed bandits, *Theor. Comput. Sci.* 410 (19) (2009) 1876–1902.
- [39] N. Merlis, S. Mannor, Batch-size independent regret bounds for the combinatorial multi-armed bandit problem, in: *Conference on Learning Theory*, PMLR, 2019, pp. 2465–2489.
- [40] X. Liu, J. Zuo, S. Wang, J.C. Lui, M. Hajiesmaili, A. Wierman, W. Chen, Contextual combinatorial bandits with probabilistically triggered arms, in: *International Conference on Machine Learning*, PMLR, 2023, pp. 22559–22593.