

TWO-STAGE GRAPH-CONSTRAINED GROUP TESTING: THEORY AND APPLICATION

Saurabh Sihag, Ali Tajer

Rensselaer Polytechnic Institute

Urbashi Mitra

University of Southern California

ABSTRACT

This paper formalizes a graph-constrained group testing (GT) framework for isolating up to k defective items from a population of p . In contrast to traditional group testing approaches, an underlying graphical model imposes constraints on how the items can be grouped for testing. The existing theories on graph-constrained GT consider one-stage, non-adaptive frameworks that can isolate the defective items perfectly with $\Theta(k^2 M^2 \log(p/k))$ tests, where M is the mixing time associated with the graph. This paper, in contrast, formalizes an *adaptive, two-stage* framework that requires $\Theta(k M^2 \log(p/k))$ tests, that is, a factor k smaller than that of the one-stage (non-adaptive) frameworks. The theoretical results established for the two-stage framework are also evaluated empirically. Furthermore, this framework is extended to address the problem of anomaly detection in the network, where based on the samples from probability distributions conforming to a location-scale family, the decision rules for detecting a defective vertex are characterized.

Index Terms— Adaptive, graph constraints, group testing.

1. INTRODUCTION

Group testing conducts pooled tests to identify defective items in a population [1]. Compared with testing the items individually, group testing approaches can provide substantial savings in the number of tests for isolating the defective items. For instance, using a non-adaptive group testing framework, $k \log \frac{p}{k}$ tests are sufficient to isolate k defectives from p items under a vanishing error criterion [2]. Due to their scalability and significant savings in cost and latency, group testing algorithms have been studied in a wide range of domains, including healthcare [3], sensor fault diagnosis [4–6], and recently, for rapid and scalable testing procedures for COVID-19 [7].

In conventional group testing approaches, selecting the items to be pooled for a test is unrestricted. However, in practice, such a selection can face constraints imposed by inherent context-dependent relationships, rendering restrictions on pooling. A graphical model can be used to model the underlying relationships and properly capture the context-dependent restrictions or preferences in pooling. For instance, in the context of infection spread in a population, it may be preferred to pool samples from members known to have more

frequent inter-personal interactions (such as family members and co-workers) over pooling arbitrarily [8]. A different relevant application of graph-constrained group testing is in network tomography, where the patterns of data traffic conform to the network connectivity, rendering constraints on pooling and are used to isolate congestion in the network [9]. In such settings, the constraints on pooling can be represented by graphical models [10].

A group testing algorithm's performance is characterized by the number of pooled tests sufficient for correctly isolating the defective items [2, 11]. When the design and outcome of a test is informed by the outcomes of previously conducted tests, the group testing framework is *adaptive* [12], and otherwise, it is *non-adaptive* [13]. In addition to the fully adaptive and non-adaptive frameworks, multi-stage testing frameworks incorporate *limited* adaptivity in group testing. In such frameworks, the tests are carried out in multiple stages. In each stage, the tests are non-adaptive, while each stage's outcome informs the design of the subsequent ones, rendering adaptivity across the stages. Such limited adaptivity can translate into significant gains over non-adaptive group testing [14–16]. An example of a multi-stage approach is a simple two-stage approach, in which the first stage is non-adaptive and the second stage depends on the first stage's outcome and consists of testing items isolated in the first stage individually [14, 15]. The interplay between the number of stages and the number of tests is investigated in [17] for two-, three-, and four-stage procedures. Besides adaptivity, group testing frameworks can also be broadly categorized as *noiseless* [11] and *noisy* [13]. In the noisy settings, tests can have erroneous outcomes.

In this paper, we formalize a two-stage graph-constrained group testing framework for isolating k defectives from a population of p items. Similar to the single-stage approach of [10], our design relies on a random walk for pooling, with the distinction that we use a *more general* random walk instead of the uniform random walk of [10]. Our analysis reveals that with M as the random walk's mixing time, we can isolate up to k defective objects with $\Theta(k M^2 \log(p/k))$ tests under the zero-error criterion. This indicates a factor of k improvement over the existing one-stage, non-adaptive, graph-constrained group testing framework [10]. Furthermore, our results show that the number of tests with two-stage group testing is approximately linear in k , which is the optimal scaling rate [1]. As an application, we also investigate applying our framework to anomaly detection in networks.

Anomaly detection is a widely studied problem of interest in networks [18, 19] and we also illustrate the application of our two-stage group testing framework to network anomaly detection. Group testing approaches for detecting anomalous sensors are investigated in [4, 5, 20]. In [4], a sensor net-

The work of S. Sihag and A. Tajer was supported in part by RPI-IBM Artificial Intelligence Research Center, and the U. S. National Science Foundation under the grants CAREER ECCS-155448 and ECCS-1933107. The work of U. Mitra was supported by the grants ONR N00014-15-1-2550, NSF CCF-1817200, ARO W911NF1910269, Cisco Foundation 1980393, DOE DE-SC0021417, Swedish Research Council 2018-04359, NSF CCF-2008927, and ONR 503400-78050.

work for a dynamic linear system is considered, and heuristic tests for detecting different faults in an arbitrary group of sensors are characterized. Algorithms based on using broadcast queries to detect dead sensors in a network are studied in [5]. In [20], a distributed group testing framework is designed for leveraging an unspecified similarity criterion to detect anomalies based on the similarity between the measurements collected by neighboring vertices. In this paper, we apply the theoretical results to an anomaly detection case study.

2. NETWORK MODEL

Consider a network represented by a *weighted, undirected, and connected* graph $\mathcal{G} \triangleq (V, E)$, where $V \triangleq \{1, \dots, p\}$ is the set of vertices and E is the set of edges. We assume that the graph is d -regular, that is, all its vertices have degree d . We use the notation $(u, v) \in E$ to denote the edge connecting vertices u and v , and denote the neighbors of $u \in V$ by $\mathcal{N}_u$. The weight associated with edge (u, v) is denoted by $w_{uv} \in \mathbb{R}_+$. In practice, the weight w_{uv} reflects the attributes that define the relationship between vertices u and v , e.g., proximity in wireless sensor networks or the likelihood of interaction between two human subjects in an epidemiological context. When $(u, v) \notin E$, we set $w_{uv} = 0$. The network, however, can contain up to k defective vertices. We denote the set of defective vertices by $\mathcal{F}$.

3. TWO-STAGE GROUP TESTING FRAMEWORK

Our goal is to characterize a *graph-constrained* two-stage testing framework that enables isolating the defective vertices $\mathcal{F}$. Our framework builds on the two-stage framework in [14] by including graphical constraint on pooling.

Stage 1. In the first stage, we conduct m tests in parallel, such that in each test, up to ℓ vertices are tested. A random walk determines the pool of vertices to be tested over the graph, and the test outcome is a binary decision determining whether the pool contains a defective vertex or not. By leveraging the outcomes of the m tests performed in this stage, the objective is to identify $2k$ vertices such that they form a superset of $\mathcal{F}$.

Stage 2. In the second stage, we individually test the $2k$ isolated vertices to identify the set $\mathcal{F}$.

Therefore, the number of tests for localizing the defective vertices is

$$\tilde{m} \triangleq m + 2k. \quad (1)$$

The group testing framework's efficiency in terms of the number of tests to be conducted hinges on the number of pooled tests in the first stage, i.e., m . In order to keep m small, as our analysis will show, the set of m random walks should collectively reach all parts of the graph reasonably fast. Next, we provide the key definitions and principles that characterize the random walk mechanism.

3.1. Test Structure

We start by formalizing a generic random walk process over the graph that furnishes the definitions and notations used in the analysis of group testing mechanism.

3.1.1. Random Walk Design

The set of vertices to be sampled in any pooled test in the first stage is determined by the vertices visited by a random walk of length ℓ . The origin of the random walk is selected uniformly at random. Any vertex may be visited more than once during a random walk, and the random walks of different tests are allowed to cross paths. We define Φ as a $p \times p$ transition matrix that models a sequential random walk on $\mathcal{G}$. The value of Φ at coordinate (i, j) , denoted by $[\Phi]_{ij}$, is given by $[\Phi]_{ij} = \phi_{ij}$, where ϕ_{ij} is the probability that the random walk transitions from the current vertex i to vertex j , when $(i, j) \in E$. We also define

$$\phi_{\max} \triangleq \max_{(i,j) \in E} \phi_{ij}, \quad \phi_{\min} \triangleq \min_{(i,j) \in E} \phi_{ij}, \quad R \triangleq \frac{\phi_{\max}}{\phi_{\min}}. \quad (2)$$

Finally, we assume that the random walk is positive recurrent and can be initialized randomly. Define $\nu \triangleq [\nu_1, \dots, \nu_p]$ as the stationary probability distribution of the random walk where ν_i is the probability that the random walk is at vertex $i \in V$ when initialized according to the stationary distribution. We also define $\lambda \triangleq \min_{u \in V} \nu_u$. The mixing time M quantifies the length of time after which the distribution of the vertices visited by any random walk on $\mathcal{G}$ becomes pointwise close to the stationary distribution ν , and it is formally defined in Definition 1.

Definition 1 (β -mixing time). *Consider a random walk W that starts at vertex $u \in V$ and terminates at time $\tau \in \mathbb{N}$. The distribution of the terminating vertex is given by ν_u^τ . Then, the β -mixing time M is the smallest integer t , such that for all $\tau \geq t$ we have*

$$\|\nu_u^\tau - \nu\|_\infty \leq \beta. \quad (3)$$

The set of vertices visited by m random walks can be formalized by a Boolean test matrix defined as follows.

Definition 2 (Test matrix). *For performing m pooled tests on a graph with p vertices, we define T as a binary matrix of dimension $m \times p$ whose elements are set according to:*

$$[T]_{tv} = \begin{cases} 1, & \text{if vertex } v \text{ is included in test } t \\ 0, & \text{otherwise} \end{cases}. \quad (4)$$

Clearly, the non-zero elements in a row t of T represent the set of vertices pooled in test t . We define $V_t \subseteq V$ as the set of vertices sampled during test t and V_t constitutes an edge-bounded and connected subgraph of $\mathcal{G}$, in which $|V_t| \leq \ell$. We can isolate the set of $2k$ vertices for individual testing in the second stage if the test matrix T satisfies certain properties [14].

3.1.2. Test Matrix Construction

We aim to assess the number of independent tests m such that T satisfies the required properties that are essential for successfully isolating $\mathcal{F}$ in a two-stage testing framework. The following definition is instrumental for characterizing the successful isolation of defective vertices.

Definition 3. A Boolean matrix T of size $m \times p$ is an (a, b, p) -selector matrix for integers $1 \leq b \leq a < p$ if any submatrix of T obtained by choosing a out of p arbitrary columns of T contains at least b distinct rows of the identity matrix I_a .

We note that for the two-stage group testing framework, the condition on T being a $(2k, k+1, p)$ -selector matrix is sufficient to isolate $\mathcal{F}$ perfectly [14]. Existing studies on graph-constrained group testing that use a random walk design, leverage the fact that matrix T being k -disjunct is a sufficient condition for isolating up to k defectives using one-stage, non-adaptive group testing [21]. In this context, we remark that a $(k+1, k+1, p)$ -selector matrix is equivalent to a k -disjunct matrix [22].

3.1.3. Decision Rules

For $t \in \{1, \dots, m\}$, the outcome of test t is binary and it flags the set V_t as anomalous if it contains at least one defective vertex. The outputs of the tests are formalized as follows.

Definition 4 (Decision rules). Test $t \in \{1, \dots, m\}$ pools the measurements from the vertices in V_t , and forms a binary decision, denoted by Δ_t , according to

$$\Delta_t = \begin{cases} 1 & \text{if the test decides } \mathcal{F} \cap V_t \neq \emptyset \\ 0 & \text{if the test decides } \mathcal{F} \cap V_t = \emptyset \end{cases}. \quad (5)$$

Accordingly, we also define $\Delta \triangleq [\Delta_1, \dots, \Delta_m]^\top$ as the decision vector.

In this paper, we focus on *noiseless* group testing where the defective vertices can be isolated correctly when the decisions in Δ are error-free. The performance of a pooled test is quantified by the probability that the outcome Δ_t is erroneous. For this purpose, we define $N(\Delta)$ as the number of erroneous decisions in Δ . The likelihood of incorrect outcomes in Δ is context-dependent and we will specify the structure of pooled tests for anomaly detection in Section 4. Therefore, in general, the accuracy of the group testing design is assessed by the likelihood that m random walks render a test matrix T that is a selector matrix with appropriate parameters and the sample complexity of pooled tests.

3.2. Two-stage Group Testing Algorithm

The two stages of group testing for successful isolation of defective vertices are described below.

Stage 1. We conduct m parallel pooled tests that are encoded by a test matrix T . If T is a $(2k, k+1, p)$ -selector matrix, the outcomes of the tests and T are jointly decoded to identify a superset of $\mathcal{F}$ that consists of $2k$ vertices.

Stage 2. Each of the $2k$ selected vertices are tested individually, leading to the identification of defective vertices $\mathcal{F}$.

The efficiency of the framework is primarily determined by Stage 1, i.e., the number of tests, m , for selector matrix construction, and the number of samples, n , for error-free outcomes in Δ with a high likelihood. These aspects are represented compactly by a binary random variable $\mathcal{P}$ defined as

$$\mathcal{P} \triangleq \begin{cases} 1 & T \text{ is } (2k, k+1, p)\text{-selector} \ \& \ N(\Delta) = 0 \\ 0 & \text{otherwise} \end{cases}$$

Therefore, the objective is to appropriately select m and n , such that, we have $\mathbb{P}(\mathcal{P} = 1) \geq 1 - \epsilon$ for some $\epsilon \in [0, 1/2]$. Next, we provide the conditions on the test matrix design and m that enable T to be a selector matrix. Subsequently, we also compare the total number of tests $\tilde{m}$ with the known information-theoretic lower bound to evaluate the performance of the two-stage framework.

Theorem 1 (Selector Matrix). When the mixing time constant β for M is set as $\beta = \frac{\lambda}{p}$, and if the parameters satisfy

$$\ell = O\left(\frac{1}{k\lambda M}\right), \quad R = O\left(\frac{p}{kM^2}\right),$$

$$d = \Omega(kM^2R), \quad m = \Theta\left(M^2\left(k\log\frac{p}{k} + \log\frac{1}{\epsilon}\right)\right),$$

then T is a $(2k, k+1, p)$ -selector matrix with a probability of at least $1 - \epsilon > 0$.

From Theorem 1, we note that the number of tests, m , for constructing the selector matrix with high probability scales as $\Theta(kM^2 \log(p/k))$. Therefore, the total number of tests $\tilde{m}$ scales as $\Theta(kM^2 \log(p/k) + k)$, which is smaller than that for one-stage group testing in [10] (approximately) by a factor k . We also note that m has quadratic dependence on the mixing time M of the random walk, which can be the dominating factor if the graph is loosely-connected. Therefore, a desirable property for the graph-constrained random walk is being able to rapidly mix (i.e., M being logarithmic in p) to avoid bottlenecks in the coverage of vertices [23]. We also remark that for a fully connected network ($M=1$), the number of tests $\tilde{m}$ achieve the information-theoretic lower bound of $k \log p/k$ in the asymptote of large p [1].

4. APPLICATION: ANOMALY DETECTION

In this section, we discuss the application of the group testing framework in Section 3 to anomaly detection. For this purpose, we start by formalizing the data model. We assume that each vertex is equipped with a sensing unit and formally, a vertex u collects a set of n identically distributed measurements denoted by $\mathbf{Y}_u^n \triangleq [Y_u^1, \dots, Y_u^n]$. When a subset of vertices $U \subseteq V$ are selected to be pooled, we denote the set of measurements collected by $\mathbf{Y}(U, n)$. When the set U does not contain any defective items, we denote the joint pdf of $\mathbf{Y}(U, n)$ by f_U^n . If vertex u is affected, the distribution of Y_u^i alters to an alternative distribution. If the set of vertices U contains at least one defective vertex, we denote the distribution of the set of measurements $\mathbf{Y}(U, n)$ by g_U^n . Finally, we define ϵ_U as the prior likelihood that the set of vertices U consists of at least one defective vertex. Furthermore, we assume that the pdfs f_U^n and g_U^n belong to the same location-scale family with equal scaling parameter and infinite support [24]. We next provide the definitions and notations to characterize the performance of a pooled test.

Pooled Test Structure: Note that the rule Δ_t conforms to the rule for the following hypothesis test:

$$\begin{aligned} H_t^0 : \quad & \mathbf{Y}(V_t, n) \sim f_{V_t}^n \\ H_t^1 : \quad & \mathbf{Y}(V_t, n) \sim g_{V_t}^n \end{aligned} \quad (6)$$

¹In practice, the distribution g_U^n is a mixture of pdfs corresponding to different realizations of defective vertices in set U .

To formalize the likelihood of erroneous decisions, we denote the true hypothesis by $T \in \{H_t^0, H_t^1\}$, and the decision by $D \in \{H_t^0, H_t^1\}$. We also define the decision vector $\delta_t(\mathbf{Y}(V_t, n)) \triangleq [\delta_t^0(\mathbf{Y}(V_t, n)), \delta_t^1(\mathbf{Y}(V_t, n))]$ for the test conducted on the set of vertices V_t with n measurements per vertex, where $\delta_t^i(\mathbf{Y}(V_t, n)) \triangleq \mathbb{1}_{\{\mathbf{Y}(V_t, n): D=H_t^i\}}$. We define $P_e(\delta_t)$ as the aggregate probability of making an incorrect decision on the true model of $\mathbf{Y}(V_t, n)$, i.e.,

$$P_e(\delta_t) = \mathbb{P}(T = H_t^0) \mathbb{P}(D = H_t^1 | T = H_t^0) + \mathbb{P}(T = H_t^1) \mathbb{P}(D = H_t^0 | T = H_t^1). \quad (7)$$

Accordingly, we denote the associated error exponent as the number of samples n grows by $\psi_t(\delta_t) \triangleq -\lim_{n \rightarrow \infty} \frac{\log P_e(\delta_t)}{n}$. We also define Ψ as the minimum error exponent among the error exponents of the m tests in test matrix T , i.e.,

$$\Psi(T) \triangleq \min_{t \in \{1, \dots, m\}} \psi_t(\delta_t). \quad (8)$$

Performance Guarantees: The following theorem specifies the decision rule that minimizes $P_e(\delta_t)$, its corresponding error exponent, and the sufficient number of samples n to ensure $N(\Delta) = 0$ with high likelihood. For this purpose, we denote the Chernoff information between two pdfs f and g by $C(f, g) \triangleq \min_{\alpha \in (0,1)} \int f^\alpha(x) g^{1-\alpha}(x) dx$.

Theorem 2. The decision rule δ_t that minimizes $P_e(\delta_t)$ for a pooled test on a set of vertices V_t is the maximum-a-posteriori (MAP) rule, and it is given by

$$\delta_t^i(\mathbf{Y}(V_t, n)) = \begin{cases} 1 & \text{if } \frac{g_{V_t}^n(\mathbf{Y}(V_t, n))}{f_{V_t}^n(\mathbf{Y}(V_t, n))} > \frac{\epsilon_{V_t}}{1 - \epsilon_{V_t}} \\ 0 & \text{otherwise} \end{cases}. \quad (9)$$

The error exponent for δ_t is given by $\psi_t(\delta_t) = C(f_{V_t}^n, g_{V_t}^n)$. Furthermore, for $\epsilon \in (0, 1/2)$, we have $\mathbb{P}(N(\Delta) > 0) \leq \epsilon$ when

$$n = \Omega\left(\frac{1}{\Psi(T)} \log \frac{m}{\epsilon}\right), \quad (10)$$

where m is determined according to Theorem 1.

The proof of $\psi_t(\delta_t)$ follows from the Chernoff-Stein Lemma for binary Bayesian hypothesis testing [25]. Therefore, the number of pooled tests, m , determined from Theorem 1 and number of samples, n , from Theorem 2 ensure the success of two-stage group testing framework for anomaly detection with high likelihood.

5. NUMERICAL RESULTS

We randomly generate testing matrices according to the conditions in Theorem 1 and evaluate the empirical likelihood of the testing matrix being a $(2k, k+1, p)$ -selector matrix. For comparison with the one-stage framework, we perform similar experiments for a $(k+1, k+1, p)$ -selector matrix. Figures 1 and 2 plot the success rate of forming a $(2k, k+1, p)$ -selector matrix and a $(k+1, k+1, p)$ -selector matrix, respectively, over 1000 random realizations. For our results,

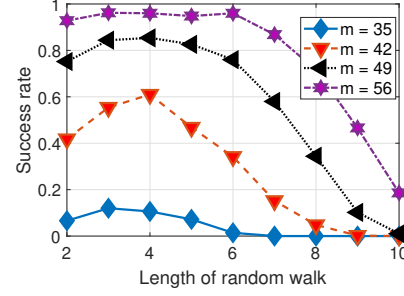


Fig. 1. Empirical likelihood of forming a $(2k, k+1, p)$ -selector matrix versus the length of random walk.

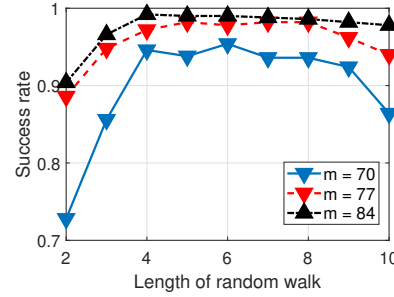


Fig. 2. Empirical likelihood of forming a $(k+1, k+1, p)$ -selector matrix versus the length of random walk.

we set $p = 100$, $k = 2$, $d = 40$. We notice that in both cases, the likelihood of forming the test matrix with the desired properties increases up to a certain length followed by a sharp decline for $(2k, k+1, p)$ -selector matrix as the length is further increased, whereas, $(k+1, k+1, p)$ -selector matrix construction is comparatively more robust to variation in the length of the random walk. This indicates that there is an upper limit on ℓ beyond which the likelihood of successful selector matrix construction diminishes, which is consistent with the implications of Theorem 1. Comparing Figures 1 and 2 reveals that constructing a $(k+1, k+1, p)$ -selector matrix requires a substantially larger number of tests compared with a $(2k, k+1, p)$ -selector matrix to achieve a similar success rate. This observation illustrates the gain offered by adopting a two-stage approach over a one-stage approach.

6. CONCLUSIONS

We have formalized a two-stage adaptive approach for graph-constrained group testing. Motivated by practical scenarios, the test design has been characterized by constraints imposed by a graphical model. Asymptotically optimal sufficient conditions on the design parameters have been established. The main observation is that introducing limited adaptivity improves the number of tests by a factor k , that is, the upper bound on the number of defective items that we aim to isolate. We have also investigated the application of this group-testing framework to anomaly detection over networks and have characterized the optimal pooled test and its sample complexity.

7. REFERENCES

- [1] Matthew Aldridge, Oliver Johnson, and Jonathan Scarlett, "Group testing: An information theory perspective," *Foundations and Trends in Communications and Information Theory*, vol. 15, no. 3-4, pp. 196–392, February 2019.
- [2] Dingzhu Du, Frank K Hwang, and Frank Hwang, *Combinatorial Group Testing and its Applications*, World Scientific, Singapore, 2000.
- [3] Ngoc T Nguyen, Hrayr Aprahamian, Ebru K Bish, and Douglas R Bish, "A methodology for deriving the sensitivity of pooled testing, based on viral load progression and pooling dilution," *Journal of Translational Medicine*, vol. 17, no. 1, pp. 252, August 2019.
- [4] Chun Lo, Mingyan Liu, Jerome P Lynch, and Anna C Gilbert, "Efficient sensor fault detection using combinatorial group testing," in *Proc. International Conference on Distributed Computing in Sensor Systems*, Cambridge, MA, 2013, pp. 199–206.
- [5] Michael T Goodrich and Daniel S Hirschberg, "Efficient parallel algorithms for dead sensor diagnosis and multiple access channels," in *Proc. ACM Symposium on Parallelism in Algorithms and Architectures*, Cambridge, MA, July 2006, pp. 118–127.
- [6] Alejandro Cohen, Asaf Cohen, and Omer Gurewitz, "Secure group testing," in *Proc. IEEE International Symposium on Information Theory*, Barcelona, Spain, July 2016, pp. 1391–1395.
- [7] Claudio M Verdun, Tim Fuchs, Pavol Harar, Dennis Elbrächter, David S Fischer, Julius Berner, Philipp Grohs, Fabian J Theis, and Felix Krahmer, "Group testing for SARS-CoV-2 allows for up to 10-fold efficiency increase across realistic scenarios and testing strategies," *medRxiv:2020.04.30.20085290*, May 2020.
- [8] Andreas Deckert, Till Bärnighausena, and Nicholas NA Kyeia, "Simulation of pooled-sample analysis strategies for COVID-19 mass testing," *Bulletin of the World Health Organization*, vol. 98, pp. 590–598, July 2020.
- [9] Weiyu Xu, Enrique Mallada, and Ao Tang, "Compressive sensing over graphs," in *Proc. IEEE International Conference on Computer Communications*, Shanghai, China, April 2011, pp. 2087–2095.
- [10] Mahdi Cheraghchi, Amin Karbasi, Soheil Mohajer, and Venkatesh Saligrama, "Graph-constrained group testing," *IEEE Transactions on Information Theory*, vol. 58, no. 1, pp. 248–262, January 2012.
- [11] Matthew Aldridge, Leonardo Baldassini, and Oliver Johnson, "Group testing algorithms: Bounds and simulations," *IEEE Transactions on Information Theory*, vol. 60, no. 6, pp. 3671–3687, March 2014.
- [12] Jacqueline M Hughes-Oliver and William H Swallow, "A two-stage adaptive group-testing procedure for estimating small proportions," *Journal of the American Statistical Association*, vol. 89, no. 427, pp. 982–993, September 1994.
- [13] George K Atia and Venkatesh Saligrama, "Boolean compressed sensing and noisy group testing," *IEEE Transactions on Information Theory*, vol. 58, no. 3, pp. 1880–1901, March 2012.
- [14] Annalisa De Bonis, Leszek Gasieniec, and Ugo Vaccaro, "Optimal two-stage algorithms for group testing problems," *SIAM Journal on Computing*, vol. 34, no. 5, pp. 1253–1270, 2005.
- [15] Marc Mézard and Cristina Toninelli, "Group testing with random pools: Optimal two-stage algorithms," *IEEE Transactions on Information Theory*, vol. 57, no. 3, pp. 1736–1745, March 2011.
- [16] Hung-Lin Fu, "Group testing with multiple mutually-obscuring positives," *Lecture Notes in Computer Science*, vol. 7777, pp. 557–568, January 2013.
- [17] Peter Damaschke and Azam Sheikh Muhammad, "Randomized group testing both query-optimal and minimal adaptive," in *Proc. International Conference on Current Trends in Theory and Practice of Computer Science*, Spindleruv Mlyn, Czech Republic, January 2012, pp. 214–225.
- [18] Ali Tajer, Venugopal V Veeravalli, and H Vincent Poor, "Outlying sequence detection in large data sets: A data-driven approach," *IEEE Signal Processing Magazine*, vol. 31, no. 5, pp. 44–56, Sep. 2014.
- [19] Dhruva Kartik, Ashutosh Nayyar, and Urbashi Mitra, "Testing for anomalies: Active strategies and non-asymptotic analysis," in *Proc. IEEE International Symposium on Information Theory*, Los Angeles, CA, June 2020, pp. 1277–1282.
- [20] Tamara Tošić, Nikolaos Thomos, and Pascal Frossard, "Distributed sensor failure detection in sensor networks," *Signal Processing*, vol. 93, no. 2, pp. 399–410, February 2013.
- [21] Anthony J Macula, "Error-correcting nonadaptive group testing with de-disjunct matrices," *Discrete Applied Mathematics*, vol. 80, no. 2-3, pp. 217–222, December 1997.
- [22] Bogdan S Chlebus and Dariusz R Kowalski, "Almost optimal explicit selectors," in *Proc. International Symposium on Fundamentals of Computation Theory*, Lubeck, Germany, August 2005, Springer, pp. 270–280.
- [23] Venkatesan Guruswami, "Rapidly mixing Markov chains: A comparison of techniques (a survey)," *arXiv:1603.01512*, 2016.
- [24] George Casella and Roger L Berger, *Statistical Inference*, vol. 2, Duxbury, Pacific Grove, CA, 2002.
- [25] Thomas M Cover and Joy A. Thomas, *Elements of Information Theory*, John Wiley and Sons, 2nd edition, 2006.