

Pavlos Nikolopoulos, Student Member, IEEE | EPFL, Lausanne, Switzerland | Email: pavlos.nikolopoulos@epfl.ch

Sundara Rajan Srinivasavaradhan ^{ID}, Member, IEEE | University of California, Los Angeles, Los Angeles, CA 90095 USA | Email: sundar@ucla.edu

Christina Fragouli ^{ID}, Fellow, IEEE | University of California, Los Angeles, Los Angeles, CA 90095 USA | Email: christina.fragouli@ucla.edu

Suhas N. Diggavi ^{ID}, Fellow, IEEE | University of California, Los Angeles, Los Angeles, CA 90095 USA | Email: suhas@ee.ucla.edu

Group Testing for Community Infections

Abstract—Group testing is the technique of pooling together diagnostic samples in order to increase the efficiency of medical testing. Traditionally, works in group testing assume that the infections are i.i.d. However, contagious diseases like COVID-19 are governed by community spread and hence the infections are correlated. This survey presents an overview of recent research progress that leverages the community structure to further improve the efficiency of group testing. We show that taking into account the side-information provided by the community structure may lead to significant savings—up to 60% fewer tests compared to traditional test designs. We review lower bounds and new approaches to encoding and decoding algorithms that take into account the community structure and integrate group testing into epidemiological modeling. Finally, we also discuss a few important open questions in this space.

Introduction

Our recent experience with COVID-19 has revealed the key role of epidemiological models and testing in the fight against pandemics (e.g., [1], [2]). For any new disease or variant of the existing ones, we will always need the ability to expeditiously deploy strategies that allow efficient testing of populations and empower targeted interventions (ideally at an individual level). This, however, poses several daunting challenges as follows.

- 1) We need to test populations at an unprecedented scale.
- 2) We need to test the same populations not just once but in a continual manner (potentially on a daily basis).
- 3) We need to estimate the epidemic state of each individual in near real time and isolate only the (most probably) infected ones.
- 4) Finally, this is to be accomplished with tests that can be limited in number and variable in speed, cost, and accuracy.

What is group testing and how can it help? Group testing is a technique that can identify the infected individuals in a population with fewer tests than the ones needed to test everyone individually. Instead of testing each person individually, group testing applies *pool tests* on the top of *groups* of diagnostic samples from multiple individuals. When pooling together these samples, particular care is taken so that the testing material is not diluted during the mixing process and that the sensitivity/specificity of the tests used is not altered significantly. The key insight is that if infections are sparse, then many group-test outcomes are likely to be negative, and therefore, all individuals included in them can be deemed healthy. However, if a group test is positive, then one cannot directly tell which individual(s) included in the test are infected; additional testing or careful decoding of other test results is therefore necessary. Accordingly, group testing offers significant benefits for sparse regimes of infection. On the other hand, if infections follow a linear or mildly sublinear regime, then individual testing has been found to be optimal [3], [4].

Group testing has a rich history dating back to Dorfman in 1943, who first introduced the concept during World War II, when the U.S. military sought to identify soldiers infected with syphilis, but tests were expensive [5]. Then on, a number of variations and setups have been examined [6]–[8].

Simply stated, the typical (static) group-testing setup assumes a population of N individuals out of which a few are infected, and the goal is to design testing strategies and corresponding decoding algorithms to identify the infections from the test results. Most works revolve around proposing a particular hand-crafted test design (e.g., random Bernoulli design) coupled with a decoding strategy (e.g., definite defectives, definite nondefectives), and guarantees are provided on the number of tests required to achieve a vanishing probability of error. In addition, order-optimal results have been proved for the asymptotic regime, where the population size tends to infinity. For example, in a population of $N \rightarrow \infty$ members, if very few people (say $k < N^{1-\Omega(1)}$) are infected, one can identify them with as low as $\mathcal{O}(k \log \frac{N}{k})$ pool tests performed in multiple adaptive stages or $\mathcal{O}(k \log N)$ pool tests performed in a single, nonadaptive stage [4], [6].¹

¹ $\mathcal{O}$ and Ω notations denote, respectively, the asymptotic upper and lower bounds, as N tends to infinity.

Interestingly, group testing is being reinvented nowadays in the context of the pandemic [9]–[14], and several countries (including India, Germany, United States, and China) have already deployed preliminary group-testing strategies [15], [16]. Also, companies and schools use pool tests to regularly monitor parts of their population and then do individual tests once a pool test comes positive (which is similar to Dorfman’s approach).

Can we do better by incorporating knowledge from a known community structure and/or epidemic dynamics? Traditional work in group testing assumes independent infections. However, viral diseases among humans have an important characteristic: infections are governed by community spread and are therefore correlated. As a use case, consider an apartment building consisting of families that have practiced social distancing; clearly, there is a strong correlation on whether members of the same family are infected or not.

In this article, we argue that taking into account the community structure may lead to significant savings in terms of the number of tests required to guarantee a given identification accuracy in the static case, or to better track the state evolution in a dynamic epidemiological model. Using entropy arguments, it is easy to see that taking into account individual correlations can help: if we represent the state (infected or not) of each individual as a binary variable, the joint entropy of correlated variables can be much smaller than the sum of the individual entropies—which is exactly the penalty we pay, if correlations are ignored. As an extreme example, assume that in each family, either all or no members are infected; then clearly, it is enough to test a single member from each family.

We also argue that leveraging the community structure can enlarge the regime, where group testing offers significant benefits over individual testing. Indeed, a limitation of group testing is that it offers very few or no benefits, when k grows linearly with N [3], [6], [17]–[19]. However, taking into account the community structure allows us to identify and remove from the population large groups of infected members, thus reducing their proportion and converting a linear to a sparse regime identification. Essentially, the community structure can guide us on when to use individual, and when group testing.

Knowing the community structure is not unrealistic. Today, it is technically feasible to keep track of the community structure—several applications are already doing so [20]–[22]. So testing according to the correlations imposed by the structure seems an approach “whose time has come,” and it has indeed attracted many researchers’ attention during the past year [23]–[31]. Moreover, it is an idea that is well aligned with the need for independent grassroots testing (schools testing their students and companies their workers) where the community structure is explicit (shared classrooms and shared common spaces).

Beyond community structure. Leveraging the community structure can be viewed as an instantiation of a recent trend in the group-testing literature and of examining variations of group testing motivated by the “real-world” scenario. For instance,

graph-constrained group testing considers the case where samples cannot be pooled together arbitrarily in a group test but must conform to constraints imposed by a graph (see, for example, [32]–[35]). Sparse group testing considers models in which individuals may participate in a limited number of tests, or tests are size constrained and cannot pool more than a limited number of samples; such constraints can significantly affect the scaling laws [36]. Different models for the test outcomes and the noise have also been considered; for instance, the work in [37] proposes a test model specifically tailored to COVID testing, where the test outcomes can provide a rough estimate of the number of infected samples. Generalized group testing subsumes as special cases a variety of noiseless and noisy group-testing models in the literature, and assumes that the test outcome is positive with some probability $f(x)$, where x is the number of defectives tested in a pool, and $f(\cdot)$ is an arbitrary monotonically increasing (stochastic) test function [38]. In this work, we do not further expand on these complementary and interesting directions.

How is this article organized? Incorporating community structure in group testing is a topic that is just emerging, and there are currently more open questions than answers. Accordingly, our goal in this article is to indicate what are potential benefits, and describe what are some first ways group testing can leverage community knowledge. After giving some background in Section “Background,” we first consider a static case in see Section “Static Case,” where we assume community structure knowledge but no knowledge of prior probabilities and aim to identify, at a particular moment in time, the infected individuals. Then, we consider dynamics in Section “Dynamic Group Testing” and aim to track the evolution of a disease over time. The static and dynamic cases are closely interrelated. For instance, the static case can be viewed as identifying the “initial state” in the case of dynamic evolution. Moreover, dynamic tracking can in some cases be reduced to simple forms of the static case, as shown through an example in Section “Dynamic Group Testing.” Most of our examples herein are from [23]–[25], [30], [31], that as far as we know were the first works to use community correlations in group test designs²; yet we also point out other very interesting works in this area [26]–[29]. Finally, Section “Conclusions and Open Questions” concludes this article with discussing open questions.

Background

Traditional (Static) Group Testing

In mathematical terms, a group test indexed by τ takes as input samples from n_τ individuals, pools them together, and outputs a single value: positive if any one of the samples is infected, and negative if none is infected. More precisely, let $U_i = 1$ when individual i is infected, and 0 otherwise. Then, the group

² Independently and in parallel, the work in [26], [27] also proposed incorporating community correlations in group test decoding.

testing output takes a binary value calculated as $Y_\tau = \bigvee_{i \in D_\tau} U_i$, where $\bigvee$ stands for the OR operator (disjunction) and D_τ is the group of people participating in the test.

Group testing typically considers the following three static models for the infections inside a population of N people.

- (i) A *combinatorial priors* model, where a fixed number of infected individuals k , is selected uniformly at random among all sets of size k .
- (ii) An *i.i.d. probabilistic priors* model, where each individual is i.i.d. infected with probability p .
- (iii) A *nonidentical probabilistic priors* model, where each item i is infected independently of all others with prior probability p_i , so that the expected number of infected members is $\bar{k} = \sum_{i=1}^N p_i$ [39].

Note that (iii) admits (ii) as a special case.

In each model, of critical interest is the minimum number of group tests $T = T(N)$ needed to identify the infected members without error or with high probability. In the combinatorial model (i), since T tests allow to distinguish among 2^T combinations of test outputs, we need $T \geq \log_2 \binom{N}{k}$ to identify k randomly infected people out of N . This is known as the *counting bound* and implies that in a sparse regime, no algorithm can use less than $T = \mathcal{O}(k \log \frac{N}{k})$ tests to achieve (almost) zero-error identification [7], [40]. In the probabilistic model (ii), a similar bound has been derived for the number of tests needed on average: $T \geq N h_2(p)$, where h_2 is the binary entropy function [6]. By extension, in the probabilistic model (iii), a lower bound for the number of tests needed can be given by the entropy, i.e., $T \geq \sum_{i=1}^N h_2(p_i)$. See [39, Appendix A] for a proof.

The usual goal in static group testing is to design a testing algorithm that is able to identify all infection statuses $\mathbf{U} = (U_1, U_2, \dots, U_N)$. These algorithms can be adaptive or nonadaptive. Adaptive testing uses the outcome of previous tests to decide what tests to perform next. An example of adaptive testing is *binary splitting*, which implements a form of binary search. Nonadaptive testing constructs, in advance, a *test matrix* $G \in \{0, 1\}^{T \times N}$ where each row corresponds to one test, each column to one member, and the nonzero elements determine the set D_τ . Although adaptive testing uses less tests than nonadaptive, nonadaptive testing is often more practical as all tests can be executed in parallel.

In the next paragraphs, we provide a brief summary of state-of-the-art algorithms for all three infection models described earlier, as well as some well-known results on their performance in various asymptotic regimes. For brevity, we focus only on the noiseless-testing case, where all group tests are supposed to be accurate, i.e., their sensitivity is 100%, although the group-testing literature also extensively studies the case when test outputs are noisy.

In the combinatorial model (i) and if the number of infected people follows a sparse regime (i.e., $k = \Theta(N^\alpha)$ and $\alpha \in [0, 1)$),

adaptive group testing, and more specifically Hwang's generalized binary splitting algorithm (BSA), is order-optimal w.r.t. the counting bound. That is, it can identify all infected individuals without error using the minimum number of tests [6], [41]. The same is also true for nonadaptive group testing whenever $\alpha \in [0, 0.409)$ and if we further allow vanishing (with N) identification errors [42]. In particular, there exists a randomized test design, coupled with a decoding algorithm named spatial inference vertex cover, which can identify all infected individuals with high probability, using the minimum number of tests. However, if $\alpha \geq 0.409$, vanishing error probabilities cannot be achieved with a single nonadaptive testing stage and at least two stages are necessary to match the counting bound [42, Ths. 1.2 and 1.3].

Conversely, classic individual testing has been proved to be optimal in the linear regime ($\bar{k} = \Theta(N)$, i.e., $\alpha = 1$). In fact, if the infection rate k/N is more than 0.38, group testing does not use fewer tests than one-to-one (individual) testing unless high identification-error rates are acceptable [3], [17]–[19]. Moreover, individual testing is preferable to nonadaptive group testing the mildly sublinear regime (where $\bar{k} = \omega(\frac{N}{\log N})$) [4].

The above mentioned achievability/converse results for the combinatorial priors are directly applicable to probabilistic model (ii) of i.i.d. priors by considering $p = k/N$. In fact, Theorems 1.7 and 1.8 from [6] imply that any algorithm that attains a vanishing probability of error on the combinatorial priors, also attains a vanishing probability of error on the corresponding i.i.d. probabilistic priors.

In the probabilistic model (iii), two well-known algorithms are the adaptive laminar algorithms that need at most $2 \sum_{i=1}^N h_2(p_i) + 2\bar{k}$ tests on average, and the “Coupon collector” nonadaptive algorithm (CCA) that needs at most $T \leq 4e(1 + \delta)\bar{k} \ln N$ test to achieve an error probability no larger than $2N^{-\delta}$ whenever $p_i \leq \frac{1}{2}$ [39], [43].

Evidently, despite the thorough analysis of the static group-testing problem, prior work has focused on independent infections. This is probably not by accident, since the problem has been motivated so far by its interesting mathematical (rather practical) aspects. Group testing is a form of inference in sparsity regimes, such as compressed sensing, but with an interesting difference: all operations are in Boolean (as opposed to real-valued) algebra, which makes the problem significantly harder. However, we believe that it is the current challenges in the context of the pandemic (e.g., scale/cost of testing) and the fact that a viral disease is indeed spread according to people's interactions that naturally bring up the need for achievability and converse results in the case of correlated/community-based infections. In Section “Static Case,” we examine this new problem.

Dynamic Infection Models

To our help, epidemiological models have been developed to describe the temporal dynamics of epidemics at different levels of detail [44], [45]. Therein, a task of interest is to track and predict the state evolution both at an individual as well as a

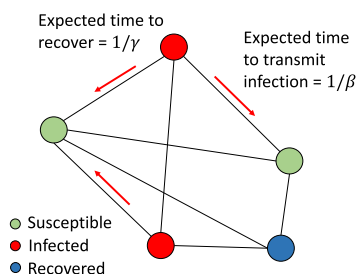


Figure 1

SIR stochastic network model. Infected nodes can potentially transmit the disease to neighboring susceptible nodes.

population level. So far, the spread of various viral diseases (including SARS-COV-2) has been studied and a number of validated models have been proposed [1], [46], [47]. The model proposed in [2] and [48] is worth mentioning as they made available an open-source library to simulate disease progression in a network while incorporating (individual-level) testing and intervention schemes.

Perhaps the most well-known model is the *continuous-time SIR stochastic network model* (see [45]), where individuals are regarded as the vertices of a graph $\mathcal{G}$ and an edge denotes a contact between neighboring vertices. At any given time, each individual can be in one of three states: susceptible, infected, or recovered. Infections can be transmitted only across an edge between a susceptible and an infected individual according to a continuous-time Markovian process of rate β (also known as transmissibility of the disease). An infected individual recovers independently of all others according to another continuous-time Markovian process of rate γ . We refer to Figure 1 for a small illustration of the SIR stochastic network model. Given any graph, this model can be exactly and efficiently simulated over a time horizon, for example, via the Gillespie algorithm (see [45, Appendix A.1]).

Static Case

In this section, we introduce community-based infections. We avoid working with the full contact graph, and prefer starting from a simplistic (yet practical) infection model to gain useful insights. This will also allow us to compute new lower bounds for the number of tests and design community-aware testing and decoding strategies. We discuss some more complex models at the end of the section.

Community-based infection model. Suppose that the total population N can be decomposed in F disjoint groups of individuals. We call these groups *families*, even though we do not refer only to actual families—we use this term to denote any group of people that happen to interact, so that they get infected according to some common principle. In addition, suppose that each family j has M_j members, so that $N = \sum_{j=1}^F M_j$.

The following community-based infection models parallel the classic ones (Section “Background”).

- **Combinatorial model (I).** k_f of the families is *infected*—namely they have at least one infected member. The rest of the families have no infected members. In each infected family j , there exist k_m^j infected members, with $0 \leq k_m^j \leq M_j$. The infected families (resp. infected family members) are chosen uniformly at random out of all families (resp. members of the same family).

- **Probabilistic model (II).** Each family is infected with probability q i.i.d. A member of an infected family j is infected, independently from the other members (and other families), with probability $p_j > 0$. If a family j is not infected, then $p_j = 0$. If $k_m^j = p_j M_j$, both models behave similarly.

Note that both these models allow families to have quite different infection levels from each other (e.g., very different infection probabilities); this is important, as, if we view the static case as a snapshot of infection evolving over time, these models enable to capture many different paths and ways to arrive at the current snapshot state.

One can remark that in model II, it seems possible that a family j is labeled “infected” without having infected members; the probability of this, however, is negligible for reasonably high infection probabilities p_j and actual values of M_j . For the purposes of this article, we will tone down this peculiarity, although the results presented further below do address it as well.

Given such community infection models, our goal is to examine whether there can be any benefits from taking the community structure into account. Some questions of interest are as follows. Is there a new lower bound on the number of tests needed for error-less identification or is the counting bound still valid? Can we design community-aware testing algorithms that are more efficient than traditional group testing, in the sense that they can achieve the same identification accuracy using significantly fewer tests? If yes, are these designs optimal and in which regimes?

Lower bounds on the number of tests. For each of the two models I and II described previously, we can compute the minimum required number of the tests using similar combinatorial and entropy arguments as in the traditional static case. We next state the lower bounds without proof; these can be found in [23, Ths. 1 and 2].

Combinatorial Community Bound

Under model I, any algorithm that identifies all k infected members without error requires a number of tests T satisfying

$$T \geq \log_2 \binom{F}{k_f} + \sum_{j=1}^{k_f} \log_2 \binom{M_j}{k_m^j}. \quad (1)$$

We can make two observations regarding the combinatorial community bound, in the case where the number of infected family members follows a “strongly” linear regime ($k_m \approx M_j$) and the number of infected families k_f follows a sparse regime (i.e., $k_f = \Theta(F^{\alpha_f})$ for $\alpha_f \in [0, 1)$).

(a) The bound increases almost *linearly* with k_f (the number of infected families), as opposed to k (the overall number of infected members). This is because, if the infection regime about families is sparse, the following asymptotic equivalence holds: $\log_2 \binom{F}{k_f} \sim k_f \log_2 \frac{F}{k_f} \sim (1 - \alpha_f) k_f \log_2 F$.

(b) In addition to the sparse regime about families, an overall sparse regime ($k = \Theta(N^\alpha)$ for $\alpha \in [0, 1)$) holds, then the community bound may be significantly lower than the counting bound that does not take into account the community structure. Consider, for example, the symmetric case, where $k_m^j = k_m$. The asymptotic behavior of the counting bound in the sparse regime is $\log_2 \binom{N}{k} \sim k \log_2 \frac{N}{k} \sim k_f k_m \log_2 \frac{F}{k_f}$, where the latter is because $k_m \approx M$. So the ratio of the counting bound to the combinatorial one scales (as F gets large) as

$$\frac{\log_2 \binom{N}{k}}{\log_2 \binom{F}{k_f} + k_f \log_2 \binom{M}{k_m}} \sim \frac{k_f k_m \log_2 \frac{F}{k_f}}{k_f \log_2 \frac{F}{k_f}} = k_m. \quad (2)$$

Although simplistic, observation (b) is important for practical reasons. Many times, the population is composed of a large number of families with members that have close contacts (e.g., relatives, work colleagues, students who attend the same classes, etc.). In such cases, we do expect that almost all members of infected families are infected (i.e., $k_m \approx M_j$), even though the overall infection regime may still be sparse. Equation (2) shows the benefits of taking the community structure into account in the test design, in such a case.

Probabilistic Community Bound

In model II, any algorithm that identifies all k infected members without error requires a number of tests T satisfying

$$T \geq F h_2(q) + \sum_{j=1}^F q M_j h_2(p_j) - w_j h_2\left(\frac{1-q}{w_j}\right) \quad (3)$$

where $w_j = 1 - q + q(1 - p_j)^{M_j}$.

Here, we make another two observations as follows.

(a) If for each family j , p_j , and M_j are such that $q(1 - p_j)^{M_j} \rightarrow 0$ (i.e., the probability of the peculiar event, where a family is labeled “infected” and yet has no infected members, is negligible), the combinatorial and probabilistic bounds are asymptotically equivalent. In particular, using the standard estimates of the binomial coefficient [49, Sec. 4.7], the combinatorial bound in (1) is asymptotically equivalent to $F h_2(k_f/F) + \sum_{j=1}^{k_f} M_j h_2(k_m^j/M_j)$, which matches the probabilistic bound in (3):

$$F h_2(q) + q \sum_{j=1}^F M_j h_2(p_j) = F h_2(\bar{k}_f/F) + \sum_{j=1}^{\bar{k}_f} M_j h_2(\bar{k}_m^j/M_j)$$

with $k_f = \bar{k}_f + o(1)$ and $k_m^j = \bar{k}_m^j + o(1)$ in place of their expected values $\bar{k}_f = Fq$ and $\bar{k}_m^j$.

(b) The probabilistic lower bound extends from zero-error recovery to constant-probability recovery by applying Fano’s inequality (as in [39, Th. 1]), and in doing so, the right-hand side (RHS) of (3) gets multiplied by the desired probability of success $\mathbb{P}(\text{suc})$.

The above mentioned results are promising, but are only possibility results. In the following, we provide example group testing algorithms that incorporate the community structure at either the encoder or decoder side.

Community structure in the test design. Our algorithm from [23], which we will simply call herein CA-adapt, is a fully adaptive community-aware test design that achieves lower bounds (1) and (3) in certain regimes.

CA-adapt consists of two parts.

The goal of the first part is to detect the infection *regime* inside each family j , so that, in the second part, the family is tested accordingly, i.e., using group testing, if j is “lightly” infected, or individual testing, otherwise. To estimate the infection regime with only a single test, the algorithm selects a random subset of representatives per each family j , r_j , it then creates a *mixed sample*³ from each subset, and finally, it applies traditional adaptive group testing on the top of all mixed samples. At the end of this step, the infection status of each mixed sample is identified.

In the second part, the algorithm treats the infection status of a mixed sample as an indicator of infection regime inside the corresponding family: if the mixed sample is positive, then the family is considered heavily infected (i.e., k_m^j/M_j or $p_j \geq 0.38$), otherwise lightly infected (i.e., k_m^j/M_j or $p_j < 0.38$). Since group testing performs better than individual testing only in the latter case (see Section “Background”), individual testing is applied to each heavily infected family member, and traditional adaptive group testing to all others.

To showcase the benefits of CA-adapt over traditional group testing, we consider the symmetric case, where $M_j = M$, $k_m^j = k_m$ (combinatorial case) or $p_j = p$ (probabilistic case), and $|r_j| = R$ for all families. Furthermore, we assume that the family representatives are selected uniformly at random without replacement, and we consider two different choices for the classic adaptive group testing used in the two parts of the algorithm: (i) Hwang’s generalized binary splitting algorithm (HGBSA) [41], which is optimal if the number of infections in the tested group is known in advance; and (ii), traditional BSA [50], which performs well, even if little is known about the number of infected members.

³ A mixed sample of a family pools together the diagnostic samples from all its representatives.

Then, in the combinatorial model I, the above mentioned community-aware algorithm succeeds using a maximum expected number of tests

$$\bar{T}_{(i)} \leq k_f \phi_c \left(\log_2 \frac{F}{k_f \phi_c} + 1 + M \right) + k(1 - \phi_c) \left(\log_2 \frac{N - k_f M \phi_c}{k(1 - \phi_c)} + 1 \right) \quad (4)$$

$$\bar{T}_{(ii)} \leq k_f \phi_c (\log_2 F + 1 + M) + k(1 - \phi_c) (\log_2 (N - k_f M \phi_c) + 1) \quad (5)$$

where the inequalities are because of the worst-case performance of HGBSA and BSA, and ϕ_c is the expected fraction of infected families whose mixed sample is positive

$$\phi_c = \begin{cases} 0, & \text{if } R = 0 \\ 1 - \frac{\binom{M-k_m}{R}}{\binom{M}{R}}, & \text{if } R \in [1, M - k_m] \\ 1, & \text{if } R \in (M - k_m, M]. \end{cases}$$

One can find analytical computations in [23], along with a similar analysis for the probabilistic infection model II. Herein, we prefer focusing on three interesting observations.

- 1) If heavily/lightly infected families are detected without errors in Part 1, CA-adapt can asymptotically achieve (up to a constant) the lower combinatorial community bound in particular cases. For example, consider a sparse regime for families (i.e., $k_f = \Theta(F^{\alpha_f})$ for $\alpha_f \in [0, 1)$) and a moderately linear regime within each family (i.e., $k_m/M \approx 0.5$). In this case,

$$\log_2 \binom{F}{k_f} \sim k_f \log_2 (F/k_f),$$

$$\log_2 \binom{M}{k_m} \sim M h_2(k_m/M) \sim M$$

and the bound in (1) becomes $k_f (\log_2 F/k_f + M)$. If R is chosen such that all infected families (which are also heavily infected as $k_m/M > 0.38$) are detected without errors (e.g., if $R > M - k_m$), then $\phi_c = 1$; thus, the RHS of (4) becomes almost equal (up to constant k_f) to the lower bound (1).

- 2) The upper bound in (5) shows that CA-adapt may achieve significant benefits compared to classic BSA in practical scenarios. If the infected families are heavily infected (as usually happens in reality) and R is chosen such that $\phi_c = 1$ (e.g., $R > M - k_m$), then $\bar{T}_{(ii)} \leq k_f (\log_2 F + 1 + M) \ll k \log_2 N + k$, where the latter is the relevant expected performance of BSA [6], [51]. Conversely, CA-adapt achieves the same performance as BSA when families are lightly infected and R is chosen such that $\phi_c = 0$ (e.g., $R = 0$). This is because $\bar{T}_{(ii)} \leq k \log_2 N + k$.
- 3) In the most favorable regime for such community-aware group testing, where very few families have almost all their members

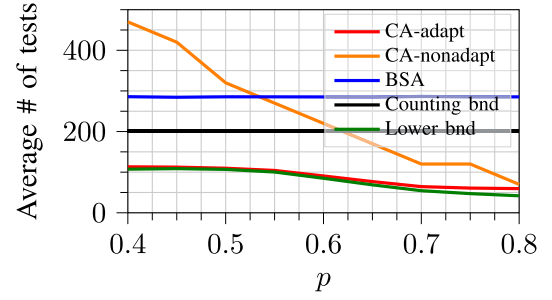


Figure 2
Tests needed by various test designs.

infected (i.e., $k_f = \Theta(F^{\alpha_f})$ for $\alpha_f \in [0, 1)$ and $k_m \approx M$), even if R is chosen optimally such that $\phi_c = 1$, the ratio of the expected number of tests needed by CA-adapt [see (4)] and HGBSA cannot be less than $1/\log(N/k)$, which upper bounds the benefits one may get. Of course, one may come up with optimized versions of CA-adapt that improve upon the gain of $1/\log(N/k)$ (see, for example, [23, Appendix B.2]).

Two remarks from the above mentioned observations are as follows. First, incorporating the community structure is more beneficial when families are heavily infected; otherwise, traditional group testing performs equally well. In fact, our experiments showed that benefits exist if the average infection rate within a family is $p > 0.15$, and increase with p . Second, a rough estimate of the families' infection rate p_j has to be known *a priori* in order to optimally choose R . As shown numerically in [23], this is unavoidable, if only a single mixed sample per family is used at the first part of CA-adapt.

In addition to adaptive community-aware group testing, there exist two-stage and nonadaptive designs [23], [27]. For example, a two-stage algorithm can be easily derived from CA-adapt, by simply replacing the adaptive algorithms (HGBSA/BSA) in both parts of the algorithm with well-known nonadaptive counterparts from the group-testing literature, such as CCW or Bernoulli designs. Following the discussion in Section "Background," such a two-stage algorithm can operate, in some regimes, with the same (order) number of tests as the adaptive algorithm, at a cost of a vanishing error probability.

Figure 2 depicts numerical evidence of how beneficial community-aware testing can be in a use case scenario of a university department with $F = 20$ classes of $M = 50$ students each, where overall infections are sparse and $p \in [0.4, 0.8]$. Two different versions of CA-adapt, one with $R = 1$ and one with $R = M$, are compared against BSA and the counting/community lower bounds. Results are averaged over 500 random communities of the same size, where infections follow model II. Interestingly, when $R = M$, CA-adapt performs close to the lower bound in most realistic scenarios $p \in [0.5, 0.8]$ (as also stated earlier). Moreover, the average overall improvement compared to traditional BSA seems to vary from 55% to 75% (fewer tests). The orange line shows a community-aware nonadaptive algorithm; even that performs better than BSA whenever $p > 0.55$ and small errors can be tolerated.

Community structure in the decoder. Consider the probabilistic model II and suppose we are interested in decoding infection status of the individuals (and families). This can be accomplished by estimating the posterior probability of the corresponding individual (or family) being infected via *loopy belief propagation* (LBP). LBP computes the posterior marginals exactly when the underlying factor graph describing the joint distribution is a tree (which is rarely the case) [52]. But, it is an algorithm of practical importance and has achieved success in a variety of applications. Also, LBP offers soft information (posterior distributions), which can be proved more useful than hard decisions in the context of disease-spread management.

We now briefly describe the factor graph and the belief propagation update rules for the probabilistic model (II). More details and exact messages can be found in [23]. Let the infection status of each family j be $V_j \sim \text{Ber}(q)$. Moreover, let $V(U_i)$ denote the family that U_i belongs to

$$\begin{aligned} &\mathbb{P}(V_1, \dots, V_F, U_1, \dots, U_N, Y_1, \dots, Y_T) \\ &= \prod_{j=1}^F \mathbb{P}(V_j) \prod_{i=1}^N \mathbb{P}(U_i | V(U_i)) \prod_{\tau=1}^T \mathbb{P}(Y_\tau | U_{\delta_\tau}) \end{aligned}$$

where δ_τ is the group of people participating in the test. The joint distribution can be represented by a factor graph, where variable nodes correspond to each random variable V_j, U_i , and Y_τ and factor nodes correspond to $\mathbb{P}(V_j)$, $\mathbb{P}(U_i | V(U_i))$, and $\mathbb{P}(Y_\tau | U_{\delta_\tau})$, respectively.

Given the result of each test is y_τ , LBP computes the marginals $\mathbb{P}(V_j = v | Y_1 = y_1, \dots, Y_T = y_T)$ and $\mathbb{P}(U_i = u | Y_1 = y_1, \dots, Y_T = y_T)$, by iteratively exchanging messages across the variable and factor nodes. The messages are viewed as *beliefs* about that variable or distributions (a local estimate of $\mathbb{P}(\text{variable} | \text{observations})$). Since all random variables are binary, each message is a 2-D vector.

We use the factor graph framework from [52] to compute the messages: variable nodes Y_τ continually transmit the message $[0, 1]$ if $Y_\tau = 1$ and $[1, 0]$ if $Y_\tau = 0$ on its incident edge, at every iteration. Each other variable node (V_j and U_i) uses the following rule: for incident each edge e , the node computes the elementwise product of the messages from every other incident edge e' and transmits this along e . For the factor node messages, we derive closed-form expressions for the sum-product update rules (akin to [52, eq. (6)]).

To showcase the benefits of community-aware decoders, we used a simple (possibly suboptimal) decoder, which is fast and can be easily configured to account for the community structure. Our LBP decoder is generic enough to accommodate any community structure and can be combined with any test design (encoder) to achieve low error rates. However, we acknowledge that many inference algorithms exist, some of which have already been employed for group testing. For example, GAMP [26], [27] and Monte Carlo sampling [53] may yield more accurate decoders.

Other community infection models. The community infection model described so far was a very simple structure with limited applications. Other works have examined more sophisticated correlations. For example, the work in [24] examines overlapping families and proposes lower bounds as well as testing and decoding strategies, based on the same principles of the two-step adaptive design and LBP decoder mentioned earlier. Numerical results indicate similar benefits; community-aware group testing needed 30%–65% fewer tests (on average) to achieve the same identification accuracy as BSA. An interesting finding was that partial knowledge of the community (e.g., knowing the families but without knowing the overlapping members) results in smaller benefits, even though it is enough to outperform community-agnostic group testing.

An even more sophisticated community model has been examined in [28], where a stochastic block model is used to describe correlated infections among families/groups.

Another approach is the linear mixing model of [27]. Contrary to the traditional Boolean formulation of the problem, that paper proposes a linear formulation with a main difference: the pooling matrix applies linear mixing to the infection statuses of the individuals, instead of disjunction operations. Because of that, the authors are able to reuse prior work on estimation with linear mixing and compressed sensing, such as the GAMP algorithm and LASSO estimator.

Dynamic Group Testing

From static to dynamic testing. So far, we considered an extension of the group testing problem that accounts for the correlated nature of infections and proposed how to modify and adapt existing techniques given the knowledge of the nature of such correlations. Taking a step back, one could further examine the problem at the source of these correlations, which is the dynamic nature of a disease. For a communicable disease, such as COVID-19, correlations are mostly induced because individuals transmit the disease continually, through direct and indirect contact. Given this fact, the assumption of a static infection model breaks down; a single round of testing is not sufficient to contain the disease—the disease continually proliferates during the testing period because of many reasons, such as inaccuracies in tests, delay in test results becoming available, insufficient testing resources, missed infections, etc. Hence, there is a need to study group testing in a nonstatic setting where one takes into account the underlying dynamics of disease progression to design tests every day, perhaps under constraints of test resources. The testing results or a summary thereof could then be used to inform individual-level decisions, such as isolation and population-scale decisions, such as a lockdown. In this section, we summarize recent progress made in this direction. The dynamic group testing problem is summarized in Figure 3 at a very high level.

Recent works have identified the significance of proactive testing and individual-level intervention for the control of the disease spread (e.g., [1], [2]). However, these solutions rely on the idea of

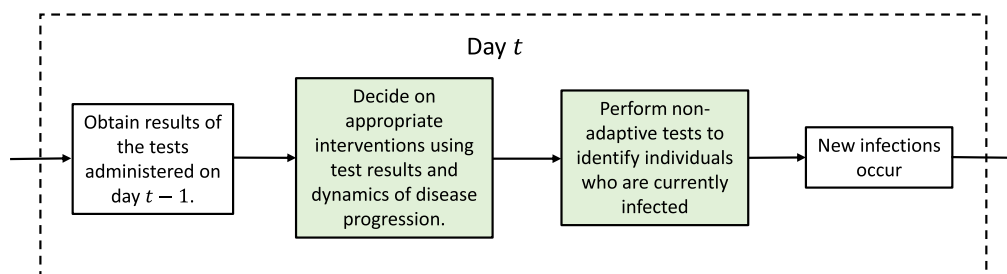


Figure 3

Dynamic testing problem with daily interventions. Test results are available 24 hours after the tests are administered. One has the flexibility to design intervention and testing modules.

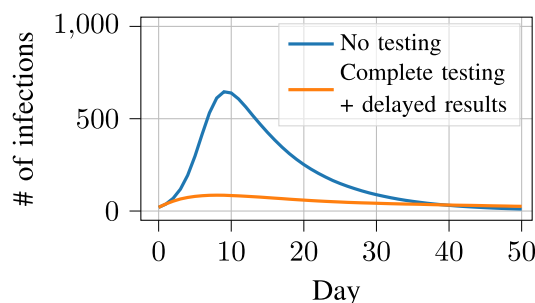


Figure 4

Simulated disease progression on a population of 1000 individuals [31]: Without any testing a large fraction of the population gets infected. Even with complete testing with delayed results (where everyone is tested individually every day, test results come out the next day, and positive cases are isolated) some infections can be missed, but “the curve is almost flattened.” So the goal of dynamic group testing is to achieve complete-testing performance using much fewer tests.

“testing everyone individually,” which can be inefficient for two reasons: on the one hand, using cheap rapid testing usually results in many people (false positives) ending up in isolation without reason and at nonnegligible societal cost; on the other hand, using accurate tests like PCR can be forbiddingly expensive. As a result, these works need to either neglect the cost of the former or alleviate the cost of the latter by scheduling tests on a (bi) weekly or monthly basis. A potential solution to the above mentioned problem is to use a small number of PCR tests every day, and exploit the power of group testing along with information from the dynamics of disease progression to inform efficient intervention schemes.

As motivated earlier, the goal of dynamic group testing is not to find all the infected individuals at a given time—in fact, this may be impossible if test resources are constrained or if test results are not instantaneous—but to contain the prevalence of the disease under the given constraints. For instance, given a fixed number of available testing resources for a week, one could intelligently choose how many tests to use and whom to test each day. Quickly identifying likely infections and isolating those individuals slows the spread of the disease (see Figure 4), which in

turn ensures that the health infrastructure is not overloaded and can operate under capacity.

Introducing pooled testing into epidemiological models. Accordingly, as a first step, the work in [30] introduces testing into epidemiological modeling and proposed dynamic testing strategies for testing every day; we summarize the ideas from this work in the next few paragraphs. We examine a simple instance of the continuous-time SIR stochastic network model (see Figure 1), where the contact $\mathcal{G}$ is a *clique*, namely a fully connected network. The focus is on a clique because of the following two reasons. First, the clique is the simplest network that one needs to understand before delving into more sophisticated ones. Second, the clique is a good model for well-mixed and closely knit communities, such as university or school classes. Our goal is to track the everyday state evolution of all individuals with the help of temporal dynamics and test results, and under the assumption that testing can happen once per day (e.g., in the morning) while its results become available only after 24 h. Further, we assume that the total number of available tests (T_{total}) over the time horizon of L days is fixed; these T_{total} tests must be distributed intelligently over these L days. For now, we ignore interventions and purely focus on state estimation. Moreover, to fix ideas, we stick to pooled testing with nonoverlapping pools.

The high-level approach to the dynamic group testing problem can be delineated as follows: steps i)–iv) repeat every day.

- (i) Obtain the test results of the previous day.
- (ii) Fix the number of tests $T^{(t)}$ to be used for the day.
- (iii) Combine the test results and the dynamics of disease progression to estimate the marginal state probabilities of each individual. This is done via estimating the posterior distribution for each individual using belief propagation and then further updating these distributions using [45, eq. (3.30)].
- (iv) Based on the estimated marginals, decide which individual goes into which test for the day.

The exact procedure for steps (ii) and (iv) is described next.

We use an *entropy reduction* approach to decide how many tests to use each day and to decide which individual goes into

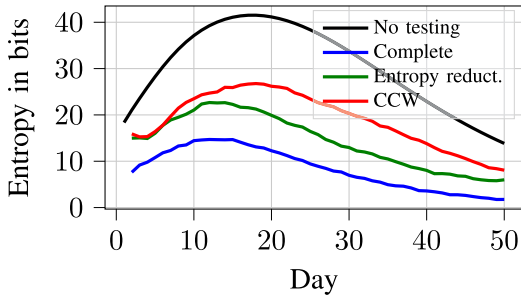


Figure 5

Uncertainty in the infection statuses of the population ($H(\mathbf{U}^{(t)})$) with different strategies. For comparison, we also plot the uncertainty when no tests are used and with complete testing, where everyone is individually tested every day.

which test. This is a common strategy employed by many adaptive algorithms in the group-testing literature to minimize the number of tests. The high-level idea in our setup is to observe that one needs more tests if the uncertainty about the state of the population is higher, and that the pools should be such that they convey maximal information about the state of the population. More precisely, suppose $\mathbf{X}^{(t)} = (X_1^{(t)}, X_2^{(t)}, \dots, X_N^{(t)})$ be the vector of SIR states of individuals at time instant t , i.e., $X_i^{(t)} \in \{S, I, R\}$. Further, let $\mathbf{U}^{(t)} = (U_1^{(t)}, U_2^{(t)}, \dots, U_N^{(t)})$ be the vector of infection statuses of the individuals at time instant t , i.e., $U_i^{(t)} = 1$ iff $X_i^{(t)} = I$ and 0 otherwise.

In order to distribute the total number of tests T_{total} over the L testing days, we do the following. On day i , say we are testing at time instant t_i . We compute $H(\mathbf{U}^{(t_i)} | \mathbf{X}^{(t_{i-1})})$ for each day i , i.e., we compute the uncertainty in the infection statuses of the population given perfect knowledge of the state of the population at the previous testing time instant. We then use this uncertainty as a guidance to decide which days needs more tests and which days need fewer tests, i.e., the number of tests used on a given day l is

$$T^{(l)} = T_{\text{total}} \frac{H(\mathbf{U}^{(t_l)} | \mathbf{X}^{(t_{l-1})})}{\sum_{i=1}^L H(\mathbf{U}^{(t_i)} | \mathbf{X}^{(t_{i-1})})}.$$

Next, we use a greedy strategy to determine the testing pools, i.e., individuals are pooled together into $T^{(l)}$ tests, such that the sum of their marginals is approximately $\frac{1}{2}$ —we do this heuristically by first placing the $T^{(l)}$ individuals whose probability of being infected are closest to $\frac{1}{2}$ in the tests, and then adding other individuals, if this moves the probability of the test result being positive closer to $\frac{1}{2}$.

Figure 5 illustrates the benefit of the entropy reduction principle compared to a static group test design CCW. With CCW, the total number of tests T_{total} was equally distributed over the testing horizon. The SIR time dynamics were simulated over a period of 50 days on a population of size 50. We observe that using the epidemic dynamics to design tests significantly reduces the uncertainty.

Introducing group testing and interventions into epidemiological models. Building on the previous idea, our work in [31] takes a theoretical approach to the dynamic group testing problem, while also incorporating interventions. Here, we are interested in the following direction: recall from Figure 4 that new infections occur daily even with complete testing; still, this is the best performance one could hope for, both in terms of containing infections and alleviating the societal impact of “false” quarantines; therefore, the question is how many tests are really needed to replicate the performance of complete testing?

We consider a discrete-time version of the SIR stochastic network model. We also consider a slightly more general graph structure than a clique—this model is termed as the *discrete-time SIR stochastic block model* as the infection model resembles a discrete-time stochastic process over a stochastic block model. More precisely, on each day each infected individual transmits the infection to a member in the same community with probability q_1 and to every other members at a lower probability q_2 . Discrete-time models fit more naturally with testing and intervention (which happen at discrete-time intervals), and are more amenable to analysis, enabling methods to derive guarantees on the number of tests needed to achieve close-to-complete-testing accuracy. The key idea is to observe that given perfect knowledge of the state of each individual the previous day, the problem reduces to that of static group testing with independent, non-identical priors.⁴

Given the above mentioned observation, we now derive the following alternate lower bound for static group testing with independent, nonidentical priors $\mathbf{p} = (p_1, p_2, \dots, p_N)$ —the number of tests needed to identify all infections is at least $\Omega(N p_{\min} \log N)$ where $p_{\min}$ is the minimum value among all p_i s. In words, this makes precise the intuition that fewer tests are necessary when infections are sparser. Moreover, if the minimum and maximum entries in $\mathbf{p}$ are of the same order, existing designs, such as CCA and CCW also require $O(N p_{\min} \log N)$ tests, and as a result, these test designs turn out to be order-optimal. We next state the conditions under which the maximum and minimum probabilities of infection: Note that if $q_1 = q_2$ (if the network is a clique), by symmetry all entries in $\mathbf{p}$ are identical and order-optimality follows. Otherwise, one could show that when q_1 and q_2 are of same order, then $p_{\max}$ and $p_{\min}$ are also of same order and order-optimality of designs, such as CCA and CCW follows. As a result, for relatively “well-mixed” populations, existing static testing strategies turn out to be order-optimal.

Simulation results show that indeed, under the above mentioned conditions on q_1 and q_2 , one could achieve the performance of complete individual testing using a much smaller number of tests via existing designs, such as CCW and CCA; for example, over a period of 50 days, group testing needs an average of around 100 tests per day for a population of 1000 individuals (see Figure 6).

⁴ Given perfect knowledge of the states 2 days prior, the problem resembles the one considered in [28].

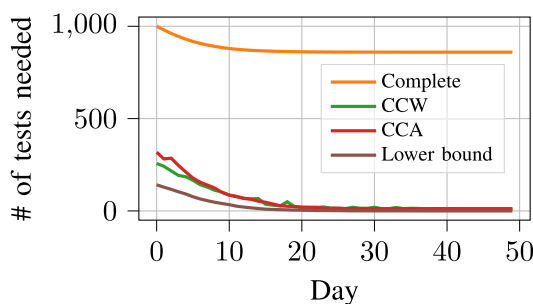


Figure 6

For the discrete-time SIR stochastic block model defined in Section “Dynamic Group Testing,” we plot the average number of tests required to identify all infected individuals every day. For comparison, we plot the entropy lower bound as well as the number of tests used by complete individual testing.

Conclusions and Open Questions

In this article, we made the case (through examples) that leveraging community structure and epidemic dynamics can enable more efficient and effective testing. But, this research direction is still largely unexplored and there exist a number of natural and important open questions, which are as follows.

Benefits. Can we gain an information-theoretic understanding over which community structures, beyond the examples we discussed, we can gain benefits, and how large these can be? A related question is, what are low complexity test and decoder designs that enable to realize such benefits. In particular, can we jointly design encoders and decoders that leverage the community structure?

Limited community structure knowledge. Limited information due to technological limitations, privacy issues, or fast-changing structures may lead to performance deterioration. This raises questions that include the following. What are “sufficient statistics” in terms of community structure information? How do community model inaccuracies affect the epidemic dynamics? How can we take into account such inaccuracies in the test designs?

Tradeoffs. In designing group testing and intervention strategies, especially over approximate dynamical models, there exist several important tradeoffs to explore, that can enable a variety of operation points. Such tradeoffs include number of tests, number of false positives (that result in unneeded isolations), number of false negatives (that can result in further disease spread), computational complexity, distributed versus centralized processing of test results, and frequency of testing.

Spectrum of tests used. Although we here focused on binary tests that are noiseless, there exists significant work on noisy testing [6], [38] as well as more involved test models [37]. Moreover, we may want to combine the use of tests that have different reliability and cost. How to optimize over a spectrum of available tests, and how the bounds and test designs would change is an open question.

Practical constraints. Deploying at large scale (e.g., over a university, or an army camp) such techniques may require us to adapt the algorithms to satisfy practical constraints; for instance, startups currently offering group testing upper limit the size of the group tested to a few tens of samples [54], which is exactly the setup examined in sparse group testing [36]. Combining constraints, such as the ones in [32], [36] with community correlations is also an open question.

Acknowledgments

This work was supported in part by NSF under Grant 2007714, Grant 1705077, and Grant 1955632. The authors would also like to thank Katerina Argyraki for her ongoing support to this project.

References

- [1] J. Taipale, P. Romer, and S. Linnarsson, “Population-scale testing can suppress the spread of COVID-19,” *medRxiv*, May 28, 2020, doi: [10.1101/2020.04.27.20078329](https://doi.org/10.1101/2020.04.27.20078329).
- [2] T. Bergstrom, C. T. Bergstrom, and H. Li, “Frequency and accuracy of proactive testing for COVID-19,” *medRxiv*, Sep. 8, 2020, doi: [10.1101/2020.09.05.20188839](https://doi.org/10.1101/2020.09.05.20188839).
- [3] M. Aldridge, “Individual testing is optimal for nonadaptive group testing in the linear regime,” *IEEE Trans. Inf. Theory*, vol. 65, no. 4, pp. 2058–2061, Apr. 2019, doi: [10.1109/TIT.2018.2873136](https://doi.org/10.1109/TIT.2018.2873136).
- [4] B. W. Heng and J. Scarlett, “Non-adaptive group testing in the linear regime: Strong converse and approximate recovery,” *Comput. Res. Repository*, 2020, *arXiv:2006.01325*. [Online]. Available: <https://arxiv.org/abs/2006.01325>
- [5] R. Dorfman, “The detection of defective members of large population,” *Ann. Math. Statist.*, vol. 14, pp. 436–440, 1943, doi: [10.1214/aoms/1177731363](https://doi.org/10.1214/aoms/1177731363).
- [6] M. Aldridge, O. Johnson, and J. Scarlett, “Group testing: An information theory perspective,” *Found. Trends Commun. Inf. Theory*, vol. 15, no. 3/4, pp. 196–392, 2019, doi: [10.1561/0100000099](https://doi.org/10.1561/0100000099).
- [7] D.-Z. Du and F. Hwang, *Combinatorial Group Testing and Its Applications* (Series on Applied Mathematics). Singapore: World Scientific, 1993, doi: [10.1142/4252](https://doi.org/10.1142/4252).
- [8] P. S. A. Yaakov Malinovsky, “Revisiting nested group testing procedures: New results, comparisons, and robustness,” *Amer. Statistician*, 2016, *arXiv:1608.06330*.
- [9] C. Gollier and O. Gossner, “Group testing against COVID-19,” *Covid Econ.*, vol. 1, no. 2, pp. 32–42, Apr. 2020. [Online]. Available: <https://hal.archives-ouvertes.fr/hal-02550740>
- [10] M. Broadfoot, “Coronavirus test shortages trigger a new strategy: Group screening,” May 2020. [Online]. Available: <https://www.scientificamerican.com/article/coronavirus-test-shortages-trigger-a-new-strategy-group-screening2/>
- [11] J. Ellenberg, “Five People. One Test. This is How You Get There,” *NYTimes*, New York, NY, USA, May 7, 2020.

- [12] C. Verdun *et al.*, "Group testing for SARS-COV-2 allows up to 10-fold efficiency increase across realistic scenarios and testing strategies," *Frontiers Public Health*, vol. 9, p. 1205, 2021, doi: [10.3389/fpubh.2021.583377](https://doi.org/10.3389/fpubh.2021.583377).
- [13] S. Ghosh *et al.*, "Tapestry: A single-round smart pooling technique for COVID-19 testing," *medRxiv*, May 2, 2020, doi: [10.1101/2020.04.23.20077727](https://doi.org/10.1101/2020.04.23.20077727).
- [14] L. M. Kucirka, S. A. Lauer, O. Laeyendecker, D. Boon, and J. Lessler, "Variation in false-negative rate of reverse transcriptase polymerase chain reaction-based SARS-COV-2 tests by time since exposure," *Ann. Intern. Med.*, vol. 173, pp. 262–267, Aug. 2020, doi: [10.7326/M20-1495](https://doi.org/10.7326/M20-1495).
- [15] S. Mallapaty, "The mathematical strategy that could transform coronavirus testing," 2020. [Online]. Available: <https://www.nature.com/articles/d41586-020-02053-6>
- [16] FDA, "Pooled sample testing and screening testing for COVID-19," 2020. [Online]. Available: <https://www.fda.gov/medical-devices/coronavirus-covid-19-and-medical-devices/pooled-sample-testing-and-screening-testing-covid-19>
- [17] L. Riccio and C. J. Colbourn, "Sharper bounds in adaptive group testing," *Taiwanese J. Math.*, vol. 4, no. 4, pp. 669–673, Dec. 2000, doi: [10.11650/twjm/1500407300](https://doi.org/10.11650/twjm/1500407300).
- [18] M. C. Hu, F. K. Hwang, and J. K. Wang, "A boundary problem for group testing," *SIAM J. Algebr. Discrete Methods*, vol. 2, pp. 81–87, 1981, doi: [10.1137/0602011](https://doi.org/10.1137/0602011).
- [19] P. Ungar, "Cutoff points in group testing," *Commun. Pure Appl. Math.*, vol. 13, pp. 49–54, 1960, doi: [10.1002/cpa.3160130105](https://doi.org/10.1002/cpa.3160130105).
- [20] C. Troncoso *et al.*, "Decentralized privacy-preserving proximity tracing," *IEEE Data Eng. Bull.*, vol. 43, pp. 36–66, 2020.
- [21] S. Azad and S. Devi, "Tracking the spread of COVID-19 in India via social networks in the early phase of the pandemic," *J. Travel Med.*, vol. 1, pp. 1–9, 2020, doi: [10.1093/jtm/taaa130](https://doi.org/10.1093/jtm/taaa130).
- [22] A. Aktay *et al.*, "Google COVID-19 community mobility reports: Anonymization process description (version 1.0)," 2020, *arXiv:2004.04145v2*.
- [23] P. Nikolopoulos, S. Rajan Srinivasavaradhan, T. Guo, C. Fragouli, and S. Diggavi, "Group testing for connected communities," in *Proc. 24th Int. Conf. Artif. Intell. Statistics*, 2021, vol. 130, pp. 2341–2349.
- [24] P. Nikolopoulos, S. R. Srinivasavaradhan, T. Guo, C. Fragouli, and S. Diggavi, "Group testing for overlapping communities," in *Proc. IEEE Int. Conf. Commun.*, 2021, pp. 1–7, doi: [10.1109/ICC42927.2021.9500791](https://doi.org/10.1109/ICC42927.2021.9500791).
- [25] P. Nikolopoulos, T. Guo, C. Fragouli, and S. Diggavi, "Community aware group testing," 2020, *arXiv:2007.08111*.
- [26] J. Zhu, K. Rivera, and D. Baron, "Noisy pooled PCR for virus testing," 2020, *arXiv:2004.02689*.
- [27] R. Goenka, S.-J. Cao, C.-W. Wong, A. Rajwade, and D. Baron, "Contact tracing enhances the efficiency of COVID-19 group testing," in *Proc. IEEE Int. Conf. Acoust., Speech Signal Process.*, 2021, pp. 8168–8172, doi: [10.1109/ICASSP39728.2021.9414034](https://doi.org/10.1109/ICASSP39728.2021.9414034).
- [28] S. Ahn, W.-N. Chen, and A. Ozgur, "Adaptive group testing on networks with community structure," 2021, *arXiv:2101.02405*.
- [29] B. Arasli and S. Ulukus, "Group testing with a graph infection spread model," 2021, *arXiv:2101.05792*.
- [30] S. R. Srinivasavaradhan, P. Nikolopoulos, C. Fragouli, and S. Diggavi, "An entropy reduction approach to continual testing," in *Proc. IEEE Int. Symp. Inf. Theory*, 2021, pp. 611–616, doi: [10.1109/ISIT45174.2021.9518188](https://doi.org/10.1109/ISIT45174.2021.9518188).
- [31] S. R. Srinivasavaradhan, P. Nikolopoulos, C. Fragouli, and S. Diggavi, "Dynamic group testing to control and monitor disease progression in a population," 2021, *arXiv:2106.10765*.
- [32] M. Cheraghchi, A. Karbasi, S. Mohajer, and V. Saligrama, "Graph-constrained group testing," *IEEE Trans. Inf. Theory*, vol. 58, no. 1, pp. 248–262, Jan. 2012, doi: [10.1109/TIT.2011.2169535](https://doi.org/10.1109/TIT.2011.2169535).
- [33] B. Spang and M. Wootters, "Unconstraining graph constrained group testing," *Leibniz Int. Proc. Informat.*, vol. 145, pp. 46:1–46:20, 2019.
- [34] A. Karbasi and M. Zadimoghaddam, "Sequential group testing with graph constraints," in *Proc. IEEE Inf. Theory Workshop*, 2012, pp. 292–296, doi: [10.1109/ITW.2012.6404678](https://doi.org/10.1109/ITW.2012.6404678).
- [35] S. Luo, Y. Matsuura, Y. Miao, and M. Shigeno, "Non-adaptive group testing on graphs with connectivity," *J. Combinatorial Optim.*, vol. 38, no. 1, pp. 278–291, 2019, doi: [10.1007/s10878-019-00379-0](https://doi.org/10.1007/s10878-019-00379-0).
- [36] V. Gandikota, E. Grigorescu, S. Jaggi, and S. Zhou, "Nearly optimal sparse group testing," *IEEE Trans. Inf. Theory*, vol. 65, no. 5, pp. 2760–2773, May 2019, doi: [10.1109/TIT.2019.2891651](https://doi.org/10.1109/TIT.2019.2891651).
- [37] R. Gabrys *et al.*, "AC-DC: Amplification curve diagnostics for COVID-19 group testing," 2021, *arXiv:2011.05223*.
- [38] J. Scarlett, "An efficient algorithm for capacity-approaching noisy adaptive group testing," in *Proc. IEEE Int. Symp. Inf. Theory*, 2019, pp. 2679–2683, doi: [10.1109/ISIT.2019.8849310](https://doi.org/10.1109/ISIT.2019.8849310).
- [39] T. Li, C. L. Chan, W. Huang, T. Kaced, and S. Jaggi, "Group testing with prior statistics," in *Proc. IEEE Int. Symp. Inf. Theory*, 2014, pp. 2346–2350, doi: [10.1109/ISIT.2014.6875253](https://doi.org/10.1109/ISIT.2014.6875253).
- [40] O. T. Johnson, "Strong converses for group testing from finite block-length results," *IEEE Trans. Inf. Theory*, vol. 63, no. 9, pp. 5923–5933, Sep. 2017, doi: [10.1109/TIT.2017.2697358](https://doi.org/10.1109/TIT.2017.2697358).
- [41] F. K. Hwang, "A method for detecting all defective members in a population by group testing," *J. Amer. Stat. Assoc.*, vol. 67, no. 339, pp. 605–608, 1972, doi: [10.2307/2284447](https://doi.org/10.2307/2284447).
- [42] A. Coja-Oghlan, O. Gebhard, M. Hahn-Klimroth, and P. Loick, "Optimal group testing," in *Proc. Mach. Learn. Res.*, Jul. 2020, vol. 125, pp. 1374–1388.
- [43] C. L. Chan, S. Jaggi, V. Saligrama, and S. Agnihotri, "Non-adaptive group testing: Explicit bounds and novel algorithms," *IEEE Trans. Inf. Theory*, vol. 60, no. 5, pp. 3019–3035, May 2014, doi: [10.1109/ISIT.2012.6283597](https://doi.org/10.1109/ISIT.2012.6283597).
- [44] N. Masuda and P. Holme, *Temporal Network Epidemiology*. New York, NY, USA: Springer, 2017.

- [45] I. Kiss, J. Miller, and P. Simon, *Mathematics of Epidemics on Networks*, vol. 46. Berlin, Germany: Springer, 2017.
- [46] A. Colubri, K. Yadav, A. Jha, and P. C. Sabeti, "Individual-level modeling of COVID-19 epidemic risk," 2020, *arXiv:2006.16761*.
- [47] S. Ubaru, L. Horesh, and G. Cohen, "Dynamic graph based epidemiological model for COVID-19 contact tracing data analysis and optimal testing prescription," 2020, *arXiv:2009.04971*.
- [48] R. McGee "Extended seirs model," 2020. [Online]. Available: <https://github.com/ryansmcgee/seirsplus/wiki/SEIRS-Model-Description>
- [49] R. Ash, *Information Theory*. New York, NY, USA: Dover, 1990.
- [50] M. Sobel and P. A. Groll, "Group testing to eliminate efficiently all defectives in a binomial sample," *Bell Syst. Tech. J.*, vol. 38, no. 5, pp. 1179–1252, 1959, doi: [10.1002/j.1538-7305.1959.tb03914.x](https://doi.org/10.1002/j.1538-7305.1959.tb03914.x).
- [51] L. Baldassini, O. Johnson, and M. Aldridge, "The capacity of adaptive group testing," in *Proc. IEEE Int. Symp. Inf. Theory*, 2013, pp. 2676–2680, doi: [10.1109/ISIT.2013.6620712](https://doi.org/10.1109/ISIT.2013.6620712).
- [52] F. R. Kschischang, B. J. Frey, and H.-A. Loeliger, "Factor graphs and the sum-product algorithm," *IEEE Trans. Inf. Theory*, vol. 47, no. 2, pp. 498–519, Feb. 2001, doi: [10.1109/18.910572](https://doi.org/10.1109/18.910572).
- [53] M. Cuturi, O. Teboul, and J.-P. Vert, "Noisy adaptive group testing using Bayesian sequential experimental design," 2020, *arXiv:2004.12508*.
- [54] Concentric by Ginkgo: Start-up offering pooled testing for k-12 schools. [Online]. Available: <https://www.concentricbyginkgo.com/faq/families/>



Pavlos Nikolopoulos (Student Member, IEEE) received the M.Sc. degree in communications and electronics from the University of Athens, Greece, in 2012, and the Ph.D. degree in computer science from École Polytechnique Fédérale de Lausanne, Switzerland, in 2018.

He is currently a Postdoctoral Researcher with the Network Architecture Lab (NAL), EPFL, and the Algorithmic Research in Network

Information Flow Lab (ARNI), University of California, Los Angeles, Los Angeles, CA, USA. Before that, he had worked for several years as a Communications Engineer with the Greek Air Force. His research interests include intersection of network security, system design and statistics, with a focus on network measurements, neutrality, and transparency. For the past year, he has worked on group testing algorithms.



Sundara Rajan Srinivasavaradhan (Member, IEEE) received the B.Tech. and M.Tech. degrees in electrical engineering from the Indian Institute of Technology, Madras, India, in 2016. He is currently working toward the Ph.D. degree in electrical engineering with the Electrical and Computer Engineering Department, University of California, Los Angeles, Los Angeles, CA, USA.

He was previously a Research Intern with Edwards Lifesciences and Microsoft Research. His research interests include statistical inference, machine learning, and discrete mathematics.



Christina Fragouli (Fellow, IEEE) received the B.S. degree from the National Technical University of Athens, Athens, Greece, and the M.Sc. and Ph.D. degrees from the University of California, Los Angeles (UCLA), Los Angeles, CA, USA, all in electrical engineering, in 1996, 1998, and 2000, respectively.

She is currently a Professor with the Electrical and Computer Engineering Department, UCLA. Her research interests include intersection of coding theory and algorithms, with applications that include network information flow, network security and privacy, wireless networks, and bioinformatics.

Dr. Fragouli was on several editorial boards and IEEE committees. She was the recipient of several awards for her work.



Suhas N. Diggavi (Fellow, IEEE) received the Ph.D. degree in electrical engineering from Stanford University, Stanford, CA, USA, in 1998. He is currently a Professor with the University of California, Los Angeles (UCLA), Los Angeles, CA, where he directs the Information Theory and Systems laboratory. His research interests include information theory and its applications to several areas including learning, security

and privacy, data compression, wireless networks, cyber-physical systems, genomics, and neuroscience.

He was the recipient of several recognitions for his research from IEEE and ACM and also faculty awards from Google, Amazon, and Facebook. He was selected as a Guggenheim Fellow in 2021. He has also organized several IEEE and ACM conferences. More information can be found at <http://licos.ee.ucla.edu>.