

Dynamic group testing to control and monitor disease progression in a population

Sundara Rajan Srinivasavaradhan*, Pavlos Nikolopoulos†, Christina Fragouli*, Suhas Diggavi*

*University of California, Los Angeles, Electrical and Computer Engineering,

email: {sundar, christina.fragouli, suhasdiggavi}@ucla.edu

†EPFL, Switzerland, email: pavlos.nikolopoulos@epfl.ch

Abstract—Proactive testing and interventions are crucial for disease containment during a pandemic until widespread vaccination is achieved. However, a key challenge remains: Can we accurately identify all new daily infections with only a fraction of tests needed compared to testing everyone, everyday? Group testing reduces the number of tests but overlooks infection dynamics and non i.i.d nature of infections in a community, while on the other hand traditional SIR (Susceptible-Infected-Recovered) models address these dynamics but don't integrate discrete-time testing and interventions. This paper bridges the gap.

We propose a “discrete-time SIR stochastic block model” that incorporates group testing and daily interventions, as a discrete counterpart to the well-known continuous-time SIR model that reflects community structure through a specific weighted graph. We analyze the model to determine the minimum number of daily group tests required to identify all infections with vanishing error probability. We find that one can leverage the knowledge of the community and the model to inform nonadaptive group testing algorithms that are order-optimal, and therefore achieve the same performance as complete testing using a much smaller number of tests.

Index Terms—Dynamic group testing, SIR stochastic network model, COVID-19 testing

I. INTRODUCTION

COVID-19 has revealed the key role of accurate epidemiological models and testing in the fight against pandemics [1]–[6]. For any new disease or variant of the existing ones, we will always need to fast develop strategies that allow efficient testing of populations and empower targeted interventions. This poses several daunting challenges: (i) we need to test populations not just once but in a continual manner (on a daily basis), and (ii) we need to estimate the epidemic state of each individual and isolate *only* the infected ones. And all these, under the accuracy and cost limitations imposed by the various types of tests.

Recent works have identified the significance of proactive testing and individual-level intervention for the control of the disease spread (e.g. [7]–[9]), but to the best of our knowledge none of them focus on test efficiency. Most solutions rely on the idea of “testing everyone individually”, which is inefficient for two reasons: on one hand, using cheap rapid testing usually results in many people (false positives) ending up in isolation without reason and at non-negligible 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.

Therefore, a critical question is still open: can we use accurate/expensive tests more efficiently? In other words, can we identify all new infections that happen each day (complete testing performance), using significantly fewer accurate/expensive tests than complete testing? Complete accurate testing (e.g. PCR) on a daily basis and isolation of infected individuals can significantly reduce the number of infected people. This is illustrated in Fig. 1 where we use the discrete-time SIR stochastic block model proposed in this paper to show how complete testing keeps the pandemic under control. Note that even with complete testing new infections still occur due to the delay between testing and receiving the test results (Fig. 1 assumes the usual delay of one day). Still, this is the best performance we can hope for, both in terms of containing infection and alleviating the societal impact of “false” quarantines; we thus ask how many tests do we really need to replicate it.

Traditional group testing strategies offer a powerful toolset for minimizing the number of tests, but they do not account for the time dynamics of a disease spread and do not take into account community structure. When the number of available tests is limited, two strategies are usually applied: sample testing, which tests only a sample of selected individuals, and/or group testing, which pools together diagnostic samples to reduce the number of tests needed to identify infected individuals in a population (e.g., see [10], [11] and references therein). Both examine a static scenario: the state of individuals is fixed (infected or not), and the goal is to identify all infected ones.

To the best of our knowledge, our work in [12] was the first paper that targeted community-aware, group-test design for the dynamic case. In that work we used the well-established continuous-time SIR stochastic network model in [13], where individuals are regarded as the vertices of a graph $\mathcal{G}$ and an edge denotes a contact between neighboring vertices, and explored group testing strategies that were able to track the epidemic state evolution at an individual level, using a small number of tests. However, due to the complexity of the continuous time model, we were not able to provide theoretical guarantees for the minimum number of tests, and although we did consider testing delays, we left interventions for future work.

TABLE I: Acronyms used throughout the paper.

Acronym	Definition
SIR	Susceptible-Infected-Recovered
CCA	Coupon Collector Algorithm
MAP	Maximum A-Posteriori
PCR	Polymerase Chain Reaction

In this paper, we allow interventions, we use discrete-time SIR models for disease spread and we derive theoretical guarantees. 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. In this paper we use a model called the “discrete-time SIR stochastic block model,” which can be considered as a discrete version of the continuous-time SIR stochastic network model over a specific type of weighted graph. The graph used captures knowledge of an underlying community structure, as discussed in Section II. In Appendix A, we compare the continuous-time model from [13] with the discrete-time model introduced in our work and justify the use of our discrete-time model. We also note that our results are applicable to a larger set of SIR models, as discussed in Section III-C.

Our main conclusion is that we can leverage the knowledge of the community and the dynamic model to inform group testing algorithms that are order-optimal and use a much smaller number of tests than complete testing to achieve the same performance. We arrive at this conclusion building on the following contributions.

We first argue that for discrete-time SIR models, given test results that identify the infection state the previous day, the problem of identifying the new infections each day reduces to static-case group testing with independent (but not identical) priors. So, existing nonadaptive algorithms such as CCA and/or random testing [14] can be reused. Figure 2 illustrates the sequence of events taking place on each day t .

We then derive a new lower bound (Theorem 2) for the number of tests needed in the case of independent (but not identical) priors. The main benefit of the new bound is not on “improving” upon the well-known entropy lower bound (stated as Lemma 3), but having a form that allows to prove order-optimality of group testing algorithms. In particular, we can prove that under mild assumptions existing nonadaptive algorithms are order-optimal in the static case (Corollary 2). This, in our opinion, is an interesting result on its own, since non-identical priors static group testing remains a relatively unexplored field compared to i.i.d. probabilistic group testing.

Finally, we derive conditions on the discrete-time SIR stochastic block model parameters under which order-optimal group test designs for the static case are also optimal for the dynamic case (Theorem 3). Simulation results show that indeed under these conditions we can achieve the performance of complete individual testing using a much smaller (close to the entropy lower bound) number of tests; 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. Our

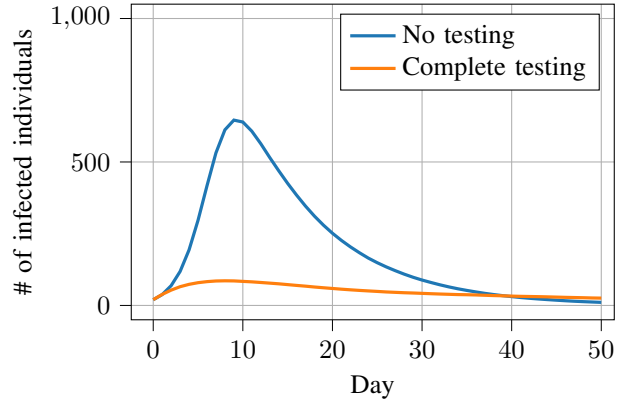


Fig. 1: Discrete-time SIR stochastic block model simulated on a population of 1000 individuals. Notice that without any testing or intervention a large fraction of the population gets infected. With *complete testing* (individually testing everyone everyday) and intervention (isolating individuals who are identified as infected) we can flatten the curve to a large extent. We assume that test results are only available the next day; if the test results were instantaneous we can identify all infections on the first day and isolate them and there would be no subsequent new infections.

simulations use existing non-adaptive test designs - we do not derive new code designs as the existing ones are sufficient. However, we do use marginal probabilities derived daily from the SIR model to inform the group design: that is, the group tests we use vary from day to day, and their design leverages the knowledge of the underlying system dynamics that depend on the community structure, as well as the previous day test results.

The rest of the paper is organized as follows: Section II provides our setup and background; Sec III contains our main results; Section IV provides numerical evaluation; Section V provides concluding discussion; Section VI discusses related work.

II. PRELIMINARIES AND PROBLEM FORMULATION

In this section we formalize our setup. Since our work for the dynamic case builds upon existing ones from static group testing, we first define the static group testing problem in Section II-A. We also review some major results in that area in Appendix B. We then provide our model (Section II-B) and problem formulation (Section II-C).

We note that throughout the paper, we assume that pooled tests are noiseless (i.e. their specificity and sensitivity rates are 100%) for simplicity. In practice, however, there might exist false positives and false negatives due to various factors, which one must consider when designing the pools. The group testing literature extensively studies the case of such noisy group testing. Additionally, for COVID-19, an estimation suggests that dilutions in the order of 5–8 would still allow a negligible false negative rate, while higher dilutions may give problems (see e.g. [15], [16] and references therein). Inevitably, all these pose practical constraints on the test matrix G .

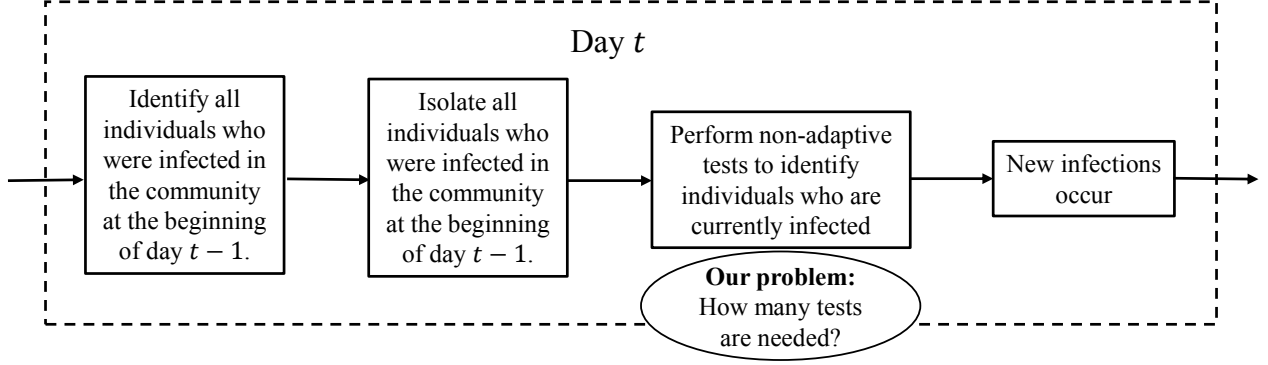


Fig. 2: The dynamic testing problem with daily interventions. How many tests are needed to achieve complete testing performance everyday, given that test results become available after a day’s delay.

A. Static group testing

Static group testing assumes a population of N individuals out of which some are infected. Each item i is infected independently of all others with prior probability p_i . A group test 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 testing output y_τ takes a binary value calculated as $y_\tau = \bigvee_{i \in \mathcal{D}_\tau} U_i$, where $\bigvee$ stands for the OR operator (disjunction) and $\mathcal{D}_\tau$ is the group of people participating in the test.

The usual goal in static group testing is to design a testing algorithm that is able to identify all infection statuses $\mathbf{U} = (U_1, \dots, U_N)$. These algorithms can be adaptive or non-adaptive. Adaptive testing uses the outcome of previous tests to decide what tests to perform next. Non-adaptive 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 non-zero elements determine the set $\mathcal{D}_\tau$. Although adaptive testing uses less tests than non-adaptive, non-adaptive testing is often more practical as all tests can be executed in parallel. We refer the reader to Appendix B for a review of results on static group testing, some of which is also used in this work.

B. Discrete-time SIR stochastic block model

We now describe our infection model via the discrete-time SIR stochastic block model with parameters $(N, C, q_1, q_2, p_{\text{init}})$. Consider a population of size N that is partitioned into multiple communities of size C . For simplicity we assume that N/C is an integer. On any day $t \in \mathbb{N}$, each individual can be in one of three states: Susceptible ($\mathcal{S}$), Infected ($\mathcal{I}$) or Recovered ($\mathcal{R}$). Let $X_i^{(t)} \in \{\mathcal{S}, \mathcal{I}, \mathcal{R}\}$ denote the state of individual i on day t , and define the state of the system as $\mathbf{X}^{(t)} \triangleq (X_1^{(t)}, X_2^{(t)}, \dots, X_N^{(t)})$. A small number of individuals are initially infected, and all new infections occur during “transmissible contacts” between infected and susceptible individuals. Recoveries occur independent of infections.

More precisely, on day $t = 0$, every individual is i.i.d infected with probability p_{init} . The following steps repeat everyday starting at $t = 1$:

- An infected individual in some community infects a susceptible one from the same community w.p. q_1 , independently of the other infected individuals of the community.
- An infected individual in some community infects a susceptible one from another community w.p. q_2 , independently from all other infected individuals.
- An infected individual recovers independently from all other individuals w.p. r .

The discrete-time SIR stochastic block model can be envisioned as a discrete version of the well-established continuous-time SIR stochastic network model [13] on the corresponding weighted graph. It inherits the main properties from the latter; for example, infections are transmitted only from an infected to a susceptible individual and both infections and recoveries are stochastic. The main difference is that in the continuous-time one, the infections and the recoveries happen according to continuous-time Markovian process with transmissibility rate β and recovery rate γ , which means that the time until a new state transition ($\mathcal{S} \rightarrow \mathcal{I}$ or $\mathcal{I} \rightarrow \mathcal{R}$) is exponentially distributed (with mean β or γ respectively). Indeed this makes the event that an individual got infected and subsequently recovered within a single day possible in the continuous-time model, whereas this is impossible in our discrete-time model.

Learning the intra-community and inter-community structure to model infection transmissions is, we believe, also practically feasible. Close contact “community” data is often readily available; for example students in each classroom in a school could form a community, and so could workers in the same office space. We also note that community-level network models alleviate some of the privacy concerns associated with using contact tracing data which tracks the exact pairs of individuals who come in contact with each other on a daily basis.

A useful remark about our model is that the state of an individual $X_i^{(t)} \in \{\mathcal{S}, \mathcal{I}, \mathcal{R}\}$ is different from the infection state $U_i^{(t)} \in \{0, 1\}$, where 1 (resp. 0) corresponds to the “infected” (resp. “not infected”). Indeed $U_i^{(t)} = 1$ iff $X_i^{(t)} = \mathcal{I}$. This difference is important in our context, because our tests do

not distinguish between susceptible and recovered individuals. In the remainder of the paper, $X_i^{(t)}$ will be called the “SIR state” of individual i , while $U_i^{(t)}$ will be i ’s “infection status.” As a result, whether a individual is infected or not changes with the day, and thus we now consider a random variable $U_i^{(t)}$ associated with each individual that describes whether it is infected on day t ($t \in N$).

C. The dynamic testing problem formulation

As can be seen from Fig. 1, testing everyone, everyday, and isolating infected individuals helps drastically reduce the number of infections that happen on a given day. We assume that the results of a test administered on a particular day are available only the next day (as usually is the case with classic PCR testing for SARS-COV-2). We also isolate only the individuals who test positive, and we do so as soon as the test results are available. Moreover, we bring back the isolated individual into the population only when they are completely recovered. Note that in the SIR model, recovered individuals cannot get infected and play no role in transmitting the infection. Therefore, without loss of generality, we could assume that isolated individuals remain isolated for the rest of the testing period.

Given these assumptions, the question we ask is if complete testing is necessary to identify all new infections everyday, or if we can achieve the same performance as complete testing with significantly fewer number of tests. In particular, how many non-adaptive group tests are necessary and sufficient to identify all new infections (with a vanishing error probability) on each day? Our problem formulation is depicted in Fig. 2. To aid a precise mathematical formulation for the problem, we first introduce some notation.

- $I_j^{(t)}$: number of new infections in community j that occurred on day t . Note that in the set-up of Fig. 2, this number is also equal to the number of infected individuals remaining in community j after intervention has been decided for day $t+1$. The new infections which happened on day t will only be identified by the tests administered on day $t+1$, whose results are available only on day $t+2$.
- $p_j^{(t)}$: the probability of an individual in community j who was susceptible at the end of day $t-1$ getting infected on day t . Note that this is same for every such individual in community j , by symmetry of the model. Moreover, we can calculate this probability as

$$\begin{aligned} p_j^{(t)} &= 1 - \mathbb{P}(\text{individual is not infected on day } t) \\ &= 1 - (1 - q_1)^{I_j^{(t-1)}} (1 - q_2)^{\sum_{j' \neq j} I_{j'}^{(t-1)}}. \end{aligned}$$

Reduction to static group testing with non-identical probabilistic priors. Note that given $I_j^{(t-1)} \forall j$, an individual belonging to community j is infected *independently* of every other individual with probability $p_j^{(t)}$ on day t . Thus, conditioned on the infection status of all individuals on day $t-1$, the infections which happen on day t are independent (but not identically distributed). Now in our dynamic testing problem set-up, on day t we perfectly learn the infection statuses of all non-isolated individuals at the time of testing on the previous

day $t-1$. Given this information, we can exactly calculate the $p_j^{(t-1)} \forall j$ (see Fig. 3), i.e. the probability that each susceptible individual in community j was later infected because of the non-isolated infected individuals.

So, given accurate test results, the dynamic testing problem is transformed daily to the problem of static group testing with non identical probabilistic priors (model (iii) in Appendix B). Therefore, the precise question we are after is the following: given that each individual in community j is infected with probability $p_j^{(t)}$ independently of every other individual, how many tests are necessary and sufficient to learn the infection status with a vanishing probability of error? We answer this question in the next section.

III. MAIN RESULTS

In this section, we prove our main theoretical results. For brevity, we will use the terms “i.i.d. priors” and “non-identical priors” to refer to i.i.d probabilistic priors model (ii) and non identical probabilistic priors model (iii) from Appendix B, respectively. The contents of this section are ordered as follows:

- First, we provide a new lower bound on the number of tests required for the problem of static group testing with non identical probabilistic priors (Theorem 2). To prove Theorem 2, we use two intermediate results: (a) we show that any test design that “works” for given prior probabilities of infection $(p_1, p_2, \dots, p_N)$ also works for the *reduced* prior probabilities $(p'_1, p'_2, \dots, p'_N)$ where $p'_i \leq p_i \leq 0.5 \forall i$. In words, we essentially prove that group testing is easier when the infections are sparser (Theorem 1); and (b) we show the following interesting property of the optimal decoder (Lemma 2) – if the optimal decoder correctly infers all the infection statuses when a set $\mathcal{D}$ is the set of infected individuals, then it will also correctly infer all the infection statuses when $\mathcal{D}' \subset \mathcal{D}$ is the set of infected individuals.
- Second, we use simple asymptotic arguments to show that some existing group testing strategies (such as CCA [14] for non-identical priors and random testing for i.i.d priors) are order-optimal for non-identical priors (Corollary 2), when $p_{\max} = O(p_{\min})$, where $p_{\max}$ is the maximum entry in $(p_1, p_2, \dots, p_N)$ and $p_{\min}$ is the minimum entry. The order $O(\cdot)$ is the standard Big-O notation used in computer science [17].
- Finally, in Theorem 3, we bridge the gap between our dynamic testing problem formulation and the above static testing problem by showing that if $q_1 = O(q_2)$, $p_{\text{init}} \leq 0.5$ and if $q_1 \leq \frac{1-1/\sqrt{2}}{C}$ and $q_2 \leq \frac{1-1/\sqrt{2}}{N}$, then the above two conditions on the prior vector are satisfied everyday in the discrete-time SIR stochastic block model parameterized by $(N, C, p_{\text{init}}, q_1, q_2)$. As a result the existing group testing strategies discussed above are order-optimal even for the dynamic testing problem formulation considered, provided that we use a sufficient number of tests each day to identify all new infections for that day.

A. Results on static group testing with non identical priors

We first consider the problem of static group testing, in which a person is infected independently with a known prior

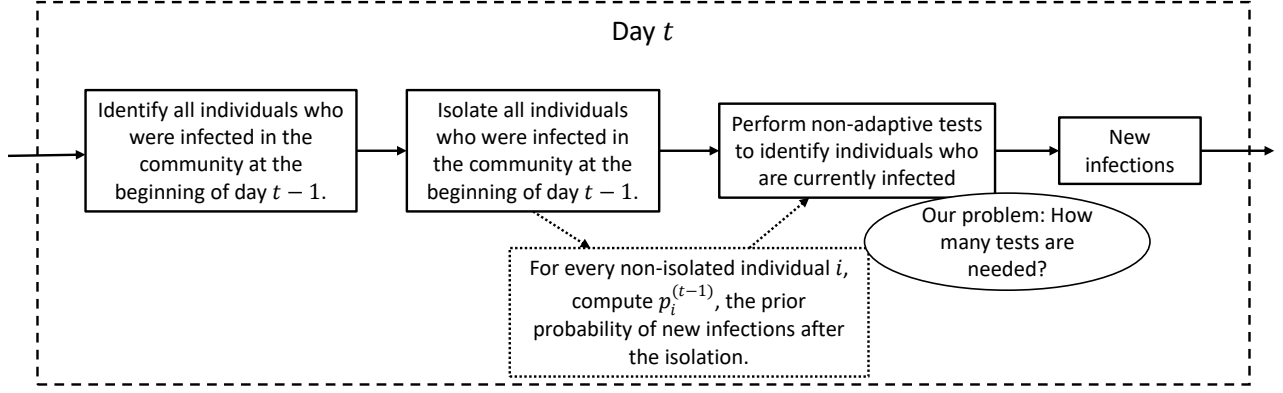


Fig. 3: From dynamic to static testing: on day t we perfectly learn the states of all non-isolated individuals at the time of testing on the previous day $t - 1$. Given this information, we know that each susceptible individual in community j is later infected with probability $p_j^{(t-1)}$ independent of every other individual. How many tests are needed to attain a vanishing probability of error on this non-identical static group testing problem?

probability p_i . Denote by $\mathbf{p} = (p_1, p_2, \dots, p_N)$ the prior vector which collects the prior probabilities of infection of all individuals. We first define some notation specific to this subsection:

- G : test matrix
- $\mathcal{D}$: set of defectives or infections
- $\mathbf{U} = (U_1, U_2, \dots, U_N)$: infection status configuration, i.e., individual i is infected if and only if $U_i = 1$.
- Define $\mathbf{U}(\mathcal{D})$ to be the vector $(U_1, \dots, U_N)$ where $U_i = 1$ iff $i \in \mathcal{D}$. Basically $\mathbf{U}(\mathcal{D})$ is the vector notation for the set of infections given by $\mathcal{D}$. For example, say $N = 3$, and $\mathcal{D} = \{1, 2\}$. Then $\mathbf{U}(\mathcal{D}) = (1, 1, 0)$. Note that there is a one-one correspondence between $\mathcal{D}$ and $\mathbf{U}(\mathcal{D})$. We will use these two notations interchangeably based on convenience.
- $G(\mathbf{U})$ represents the test results corresponding to the given test design and infection status configuration.
- For a fixed number of tests T , define a decoding function $R : \{0, 1\}^{[T]} \rightarrow \{0, 1\}^{[N]}$ which estimates the infection statuses from the test results.
- A defective set $\mathcal{D}$ “explains” test results $\mathbf{y}$ iff $G(\mathbf{U}(\mathcal{D})) = \mathbf{y}$.
- $\mathbb{P}(\mathbf{U}; \mathbf{p})$ denotes the probability of the infection status configuration under priors $\mathbf{p}$, i.e.

$$\mathbb{P}(\mathbf{U}; \mathbf{p}) = \prod_{i=1}^N p_i^{U_i} (1 - p_i)^{1-U_i}.$$

- Probability of error for a test matrix, decoder pair under given priors

$$\begin{aligned} \mathbb{P}_{err}(G, R; \mathbf{p}) &\triangleq \mathbb{E}_{\mathbf{U} \sim \mathbf{p}} \mathbb{1}\{R(G(\mathbf{U})) \neq \mathbf{U}\} \\ &= \sum_{\mathbf{u} \in \{0, 1\}^N} \mathbb{P}(\mathbf{U} = \mathbf{u}; \mathbf{p}) \mathbb{1}\{R(G(\mathbf{u})) \neq \mathbf{u}\}. \end{aligned}$$

Definition 1 (MAP decoder). For fixed priors $\mathbf{p}$ and testing matrix G with number of tests T , we define the corresponding MAP decoder as $R_{map}(\cdot; G, \mathbf{p}) : \{0, 1\}^T \rightarrow \{0, 1\}^N$, where

$$R_{map}(\mathbf{y}; G, \mathbf{p}) = \arg \max_{\mathbf{U}: G(\mathbf{U}) = \mathbf{y}} \mathbb{P}(\mathbf{U}; \mathbf{p}).$$

In case of ties, the MAP decoder will select the solution which

comes the earliest lexicographically¹.

In words, the MAP decoder chooses the most likely configuration which explains the test results. We next show that the MAP decoder is the optimal decoder for a fixed G and $\mathbf{p}$, i.e., the MAP decoder minimizes the probability of error amongst all decoders for any $G, \mathbf{p}$.

Remark. Though the MAP decoder is optimal, it is unclear if the optimization problem corresponding to the MAP decoder can be solved efficiently. However, many heuristics such as belief propagation (see for example [18]) and random sampling methods exist which approximate well the MAP decoder. That said, in this work we use the MAP decoder only as a tool for theoretical analysis of the error probability.

Lemma 1 (Optimality of MAP decoder). For given test matrix G and priors $\mathbf{p}$, the corresponding MAP decoder minimizes the probability of error for the test matrix under the given priors, i.e.,

$$\mathbb{P}_{err}(G, R_{map}(\cdot; G, \mathbf{p}); \mathbf{p}) \leq \mathbb{P}_{err}(G, R; \mathbf{p}) \quad \forall R.$$

We provide the proof of Lemma 1 in Appendix C.

Given the optimality of the MAP decoder, we will denote by $\mathbb{P}_{err}^*(G, \mathbf{p}) \triangleq \mathbb{P}_{err}(G, R_{map}(\cdot; G, \mathbf{p}); \mathbf{p})$, the optimal probability of error corresponding to a given test design and priors.

We next prove a property of the MAP decoder, and this property will be used in the proof of our main result that follows. The following Lemma says that suppose the MAP decoder correctly identifies a given defective set, then it will always correctly identify a sparser defective set.

Lemma 2. Consider a test matrix G and priors $\mathbf{p}$, and let $p_i \leq 0.5$ for all i . Suppose the corresponding MAP decoder is erroneous when identifying the defective set $\mathcal{D}$. Then the MAP decoder is also erroneous when identifying the defectives set

¹We only include this to make the definition of the MAP decoder deterministic and to simplify subsequent analysis. This particular choice has no significance and all results still hold regardless.

$\mathcal{D} \cup \{j\}$, i.e.,

$$\begin{aligned} & \mathbb{1}\{R_{\text{map}}(G(\mathbf{U}(\mathcal{D} \cup \{j\})); G, \mathbf{p}) \neq \mathbf{U}(\mathcal{D} \cup \{j\})\} \\ & \geq \mathbb{1}\{R_{\text{map}}(G(\mathbf{U}(\mathcal{D})); G, \mathbf{p}) \neq \mathbf{U}(\mathcal{D})\}. \end{aligned}$$

For the proof of Lemma 2, we refer the reader to Appendix D.

We next prove the main new result for the static case. In words, the following theorem says that the group testing problem is not harder when the infections are sparser. As a result, this allows us to lower/upper bound the number of tests needed for the group testing problem with non-identical priors by the number of tests needed for an alternative group testing problem with identical priors.

Theorem 1. *Consider a testing matrix G used with two different sets of priors $\mathbf{p}$ and $\mathbf{p}'$. For a given $j \in [N]$, where $[N]$ is defined as the set $\{1, 2, \dots, N\}$, let $p'_j \leq p_j \leq 0.5$ and let $p'_i = p_i$ for every $i \in [N], i \neq j$. The two prior vectors are same everywhere except at index j where $\mathbf{p}'$ is smaller. Then*

$$\mathbb{P}_{\text{err}}^*(G, \mathbf{p}') \leq \mathbb{P}_{\text{err}}^*(G, \mathbf{p}).$$

Proof. We prove this by showing that when the MAP decoder corresponding to $(G, \mathbf{p})$ pair is used as a decoder with $(G, \mathbf{p}')$, the probability of error is always lower, i.e.,

$$\begin{aligned} & \mathbb{P}_{\text{err}}(G, R_{\text{map}}(\cdot; G, \mathbf{p}); \mathbf{p}') \\ & \leq \mathbb{P}_{\text{err}}(G, R_{\text{map}}(\cdot; G, \mathbf{p}); \mathbf{p}) = \mathbb{P}_{\text{err}}^*(G, \mathbf{p}). \end{aligned}$$

As a result the optimal decoder for $(G, \mathbf{p}')$ pair has a probability of error not exceeding this quantity.

Now, we can express the probability of error for the MAP decoder of $(G, \mathbf{p})$ pair. For simplicity of notation in the following derivations, $\mathcal{E}(\mathcal{D}) \triangleq \mathbb{1}\{R_{\text{map}}(G(\mathbf{U}(\mathcal{D})); G, \mathbf{p}) \neq \mathbf{U}(\mathcal{D})\}$ denotes the indicator of the event that the MAP decoder is erroneous when the defective set is $\mathcal{D}$ (and under further assumptions that the priors are $\mathbf{p}$ and test matrix is G).

$$\begin{aligned} & \mathbb{P}_{\text{err}}(G, R_{\text{map}}(\cdot; G, \mathbf{p}); \mathbf{p}) \\ &= \sum_{\mathcal{D} \subseteq [N]} \mathbb{P}(\mathbf{U}(\mathcal{D})) \mathcal{E}(\mathcal{D}) \\ &= \sum_{\mathcal{D} \subseteq [N]} \prod_{i \in \mathcal{D}} p_i \prod_{l \in [N] \setminus \mathcal{D}} (1 - p_l) \mathcal{E}(\mathcal{D}) \\ &\stackrel{(a)}{=} \sum_{\mathcal{D} \subseteq [N] | j \in \mathcal{D}} \prod_{i \in \mathcal{D}} p_i \prod_{l \in [N] \setminus \mathcal{D}} (1 - p_l) \mathcal{E}(\mathcal{D}) \\ &\quad + \sum_{\mathcal{D} \subseteq [N] | j \notin \mathcal{D}} \prod_{i \in \mathcal{D}} p_i \prod_{l \in [N] \setminus \mathcal{D}} (1 - p_l) \mathcal{E}(\mathcal{D}) \\ &\stackrel{(b)}{=} p_j \sum_{\mathcal{D} \subseteq [N] \setminus \{j\}} \prod_{i \in \mathcal{D}} p_i \prod_{l \in [N] \setminus \mathcal{D} \cup \{j\}} (1 - p_l) \mathcal{E}(\mathcal{D} \cup \{j\}) \\ &\quad + (1 - p_j) \sum_{\mathcal{D} \subseteq [N] \setminus \{j\}} \prod_{i \in \mathcal{D}} p_i \prod_{l \in [N] \setminus \mathcal{D} \cup \{j\}} (1 - p_l) \mathcal{E}(\mathcal{D}), \quad (1) \end{aligned}$$

where in (a) we split the summation into two cases – one where $j \in \mathcal{D}$ and the other where $j \notin \mathcal{D}$; in (b) we take j out of the summation.

Similarly, one could express the probability of error for the

same decoder with the pair $(G, \mathbf{p}')$ as

$$\begin{aligned} & \mathbb{P}_{\text{err}}(G, R_{\text{map}}(\cdot; G, \mathbf{p}); \mathbf{p}') \\ &= p'_j \sum_{\mathcal{D} \subseteq [N] \setminus \{j\}} \prod_{i \in \mathcal{D}} p_i \prod_{l \in [N] \setminus \mathcal{D} \cup \{j\}} (1 - p_l) \mathcal{E}(\mathcal{D} \cup \{j\}) \\ &\quad + (1 - p'_j) \sum_{\mathcal{D} \subseteq [N] \setminus \{j\}} \prod_{i \in \mathcal{D}} p_i \prod_{l \in [N] \setminus \mathcal{D} \cup \{j\}} (1 - p_l) \mathcal{E}(\mathcal{D}). \quad (2) \end{aligned}$$

The first error term $\mathbb{P}_{\text{err}}(G, R_{\text{map}}(\cdot; G, \mathbf{p}); \mathbf{p})$ is of the form $p_j a + (1 - p_j) b$, and the second error term $\mathbb{P}_{\text{err}}(G, R_{\text{map}}(\cdot; G, \mathbf{p}); \mathbf{p}')$ is of the form $p'_j a + (1 - p'_j) b$. From Lemma 2, we have $\mathcal{E}(\mathcal{D} \cup \{j\}) \geq \mathcal{E}(\mathcal{D})$ and hence $a \geq b$. Since $a \geq b$ and $p'_j \leq p_j$, one can verify that $p_j a + (1 - p_j) b \geq p'_j a + (1 - p'_j) b$, and thus

$$\mathbb{P}_{\text{err}}(G, R_{\text{map}}(\cdot; G, \mathbf{p}); \mathbf{p}) \geq \mathbb{P}_{\text{err}}(G, R_{\text{map}}(\cdot; G, \mathbf{p}); \mathbf{p}'),$$

concluding the proof. $\square$

Now one could repeatedly apply Theorem 1 on the prior vector $\mathbf{p}$ to conclude that any test matrix G should only do better on the reduced uniform prior vector $\mathbf{p}_{\min} = (p_{\min}, p_{\min}, \dots, p_{\min})$ where $p_{\min} \triangleq \min_{i \in [N]} p_i$. On the other hand, the test matrix G should only do worse on the prior vector $\mathbf{p}_{\max} = (p_{\max}, p_{\max}, \dots, p_{\max})$ where $p_{\max} \triangleq \max_{i \in [N]} p_i$. This is stated below as a corollary.

Corollary 1. *Consider a test matrix G and a prior vector $\mathbf{p}$ such that $p_i \leq 0.5$ for all $i \in [T]$. Let $\mathbf{p}_{\min} = (p_{\min}, p_{\min}, \dots, p_{\min})$ where $p_{\min} \triangleq \min_{i \in [N]} p_i$ and let $\mathbf{p}_{\max} = (p_{\max}, p_{\max}, \dots, p_{\max})$ where $p_{\max} \triangleq \max_{i \in [N]} p_i$. Then*

$$\mathbb{P}_{\text{err}}^*(G, \mathbf{p}_{\max}) \geq \mathbb{P}_{\text{err}}^*(G, \mathbf{p}) \geq \mathbb{P}_{\text{err}}^*(G, \mathbf{p}_{\min}).$$

As a consequence of the above corollary, the number of tests required to attain a fixed (small) probability of error ϵ with prior vector $\mathbf{p}_{\min}$ is not more than the number of tests required to attain probability of error ϵ with prior vector $\mathbf{p}$. This observation allows us to use the lower bound on the number of tests when the priors are identical. This is made precise in the following theorem.

Theorem 2. *Consider the non-adaptive group testing problem with N items where the probability of item i being infected is $p_i \leq 0.5$. Let $p_{\min} \triangleq \min_{i \in [N]} p_i$. In order to achieve a probability of error $\rightarrow 0$ as $N \rightarrow \infty$, the number of tests must be*

$$T(\mathbf{p}) = \Omega(\min\{N, N p_{\min} \log N\}).$$

Proof. From Corollary 1, suppose a test matrix G achieves a probability of error ϵ on prior vector $\mathbf{p}$, the same test matrix achieves a probability of error not more than ϵ on the prior vector $\mathbf{p}_{\min}$, where $\mathbf{p}_{\min} = (p_{\min}, p_{\min}, \dots, p_{\min})$ and $p_{\min} \triangleq \min_{i \in [N]} p_i$. Any strategy that achieves a probability of error $\rightarrow 0$ as $N \rightarrow \infty$ with the prior vector $\mathbf{p}_{\min}$ requires a number of tests equal to $\Omega(\min\{N, N p_{\min} \log N\})$. Thus, we need at least as many tests with the prior vector $\mathbf{p}$. $\square$

As discussed in Appendix B, the entropy bound in Lemma 3 is an alternate lower bound on the number of tests needed for this problem. We note that the entropy bound might be

greater or smaller than the term $Np_{\min} \log N$ in Theorem 2. In particular, if $p_i \leq 1/2 \forall i$ it is easy to see that $\sum_{i=1}^N h_2(p_i) \geq N h_2(p_{\min}) \geq N p_{\min} \log 1/p_{\min}$. However the term $1/p_{\min}$ may be smaller or larger than N ; thus our bound, that applies independently of the value of $p_{\min}$ (as long as $p_i \leq 0.5$) cannot be directly derived from the entropy bound, and could be either greater or lesser than the entropy bound. Having said that, the main advantage of the lower bound in Theorem 2 is its particular form, which allows the proof of order-optimality of several static group testing algorithms, as we will see in the next subsection.

Now, if the prior vector $\mathbf{p}$ is “bounded”, in the sense that the maximum entry and minimum entry in $\mathbf{p}$ differ by a constant factor (constant with respect to N), then the lower bound can be re-written in terms of the maximum entry in $\mathbf{p}$ or the mean of $\mathbf{p}$. Basically we here just use the fact that constant factors do not affect the order. We next make this corollary precise.

Definition 2 (Bounded priors). *Let $\eta \in [1, \infty)$ be a fixed constant (constant with respect to N). A prior vector $\mathbf{p}$ of length N is called η -bounded if*

$$\frac{\max_i p_i}{\min_i p_i} \leq \eta.$$

Corollary 2 (Lower bound for bounded priors). *Consider the non-adaptive group testing problem with N items where the probability of item i being infected is $p_i \leq 0.5$. Let $p_{\max} \triangleq \max_{i \in [N]} p_i$ and $p_{\text{mean}} \triangleq \frac{1}{N} \sum_{i=1}^N p_i$. Suppose $\mathbf{p} = (p_1, \dots, p_N)$ is η -bounded for some constant η . Any strategy that achieves a probability of error $\rightarrow 0$ as $N \rightarrow \infty$ requires*

$$\begin{aligned} T(\mathbf{p}) &= \Omega(\min\{N, Np_{\text{mean}} \log N\}) \\ &= \Omega(\min\{N, Np_{\max} \log N\}). \end{aligned}$$

B. Performance of existing non-adaptive algorithms in the static non-identical priors

Suppose $\mathbf{p}$ is η -bounded and each $p_i \leq 0.5$. The following non-adaptive algorithms can be proved to be order-optimal with respect to the lower bound in Corollary 2:

- The Coupon Collector Algorithm (CCA) from [14] for prior vector $\mathbf{p}$, as discussed in Appendix B, achieves a probability of error less than $2N^{-\delta}$ with a number of tests less than $4e(1+\delta)Np_{\text{mean}} \log N$ (see Theorem 3 in [14]). As a result, w.r.t to the lower bound in Corollary 2, either CCA is order-optimal (if $N \geq Np_{\text{mean}} \log N$) or individual testing is order optimal (if $N \leq Np_{\text{mean}} \log N$).
- As discussed in Appendix B for the group testing problem with identical priors (say every item is infected with the same probability p'), a variety of randomized and explicit algorithms have been proposed² which achieve a vanishing probability of error with a number of tests $O(Np' \log N)$. From Corollary 1, any test matrix that achieves a vanishing probability of error with $\mathbf{p}_{\max}$ should also attain a vanishing

²Most of these were considered in the context of combinatorial priors. However, Theorem 1.7 and Theorem 1.8 from [10] 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.

probability of error with $\mathbf{p}$, and as a result $O(Np_{\max} \log N)$ tests are sufficient for the prior vector $\mathbf{p}$. Consequently w.r.t our lower bound in Corollary 2, any of these designs is order optimal (if $N \geq Np_{\max} \log N$) or individual testing is order optimal (if $N \leq Np_{\max} \log N$).

C. Dynamic testing - bridging the gap

Given the discussion above, we next show conditions under which the prior probabilities of infections each day (these change everyday) are η -bounded and are each not more than 0.5. If these two conditions are satisfied everyday for our discrete-time SIR stochastic block model set-up in Fig. 3, then CCA and the other algorithms discussed in Section III-B are order-optimal for our dynamic testing problem formulation. (see 3). We first define some notation, building upon the notation in Section II-C.

- $p_{\max}^{(t)} \triangleq \max_j p_j^{(t)}$, the maximum probability of new infection on day t .
- $p_{\min}^{(t)} \triangleq \min_j p_j^{(t)}$, the minimum probability of new infection on day t .

Theorem 3. *Consider the testing-intervention problem in Fig. 3 where the infections follow the discrete-time SIR stochastic block model $(N, C, q_1, q_2, p_{\text{init}})$.*

- Suppose $p_{\text{init}} \leq 0.5$, $q_1 \leq \frac{1-1/\sqrt{2}}{C}$ and $q_2 \leq \frac{1-1/\sqrt{2}}{N}$, then $p_j^{(t)} \leq 0.5$.*
- Suppose $\frac{q_1}{q_2} \leq \eta$, then $\frac{p_{\max}^{(t)}}{p_{\min}^{(t)}} \leq \eta$ and as a result the prior vector for each day is η -bounded.*

Proof. We prove (i) first. We first have $p_j^{(0)} = p_{\text{init}} \leq 0.5$. For $t \geq 1$, we have

$$\begin{aligned} p_j^{(t)} &= 1 - (1 - q_1)^{I_j^{(t-1)}} (1 - q_2)^{\sum_{j' \neq j} I_{j'}^{(t-1)}} \\ &\stackrel{(a)}{\leq} 1 - (1 - q_1)^C (1 - q_2)^N \\ &\stackrel{(b)}{\leq} 1 - (1 - Cq_1)(1 - Nq_2) \stackrel{(c)}{\leq} 0.5, \end{aligned}$$

where in (a) we used the fact that the total number of infections inside a community and overall cannot be greater than C and N , respectively; (b) follows because of the algebraic inequality $(1+x)^y \geq 1+xy$ if $x \geq -1$ and $y \notin (0, 1)$; in (c) we used our assumptions about q_1 and q_2 .

We next prove (ii). Since $q_1 \geq q_2$ in our model, we have

$$\begin{aligned} p_j^{(t)} &= 1 - (1 - q_1)^{I_j^{(t-1)}} (1 - q_2)^{\sum_{j' \neq j} I_{j'}^{(t-1)}} \\ &= 1 - \left(\frac{1 - q_1}{1 - q_2} \right)^{I_j^{(t-1)}} (1 - q_2)^{\sum_{j' \neq j} I_{j'}^{(t-1)}} \\ &\leq 1 - \left(\frac{1 - q_1}{1 - q_2} \right)^{\max_j I_j^{(t-1)}} (1 - q_2)^{\sum_{j' \neq j} I_{j'}^{(t-1)}}, \quad (3) \end{aligned}$$

where $\max_j I_j^{(t)}$ is simply the maximum number of infections over all communities. Likewise,

$$\begin{aligned} p_j^{(t)} &= 1 - (1 - q_1)^{I_j^{(t-1)}} (1 - q_2)^{\sum_{j' \neq j} I_{j'}^{(t-1)}} \\ &\stackrel{(a)}{\geq} 1 - (1 - q_2)^{\sum_{j'} I_{j'}^{(t-1)}}, \quad (4) \end{aligned}$$

where in (a) we used $q_2 \leq q_1$. Combining (3) and (4) we have

$$\begin{aligned} \frac{p_{\max}^{(t)}}{p_{\min}^{(t)}} &= \frac{1 - \left(\frac{1-q_1}{1-q_2}\right)^{\max_j I_j^{(t-1)}} (1-q_2)^{\sum_{j'} I_{j'}^{(t-1)}}}{1 - (1-q_2)^{\sum_{j'} I_{j'}^{(t-1)}}} \\ &\stackrel{(a)}{\leq} \frac{1 - \left(\frac{1-q_1}{1-q_2}\right)^{\max_j I_j^{(t-1)}} (1-q_2)^{\max_j I_j^{(t-1)}}}{1 - (1-q_2)^{\max_j I_j^{(t-1)}}} \\ &= \frac{1 - (1-q_1)^{\max_j I_j^{(t-1)}}}{1 - (1-q_2)^{\max_j I_j^{(t-1)}}} \stackrel{(b)}{\leq} \frac{q_1}{q_2} \leq \eta. \end{aligned}$$

where (a) follows from the following facts: the function $f_1(x) = \frac{1-\kappa x}{1-x}$ is increasing for $\kappa \in (0, 1)$, the function $f_2(x) = (1-q_2)^x$ is decreasing for $q_2 \in (0, 1)$, and the sum $\sum_{j'} I_{j'}^{(t-1)}$ is lower bounded by $\max_j I_j^{(t-1)}$; and (b) follows from the fact that the function $f_3(x) = \frac{1-(1-q_1)x}{1-(1-q_2)x}$ is decreasing in $x \geq 1$ for $q_1 \geq q_2$, and therefore the maximum of the ratio is obtained for $\max_j I_j^{(t-1)} = 1$. All proofs of the above statements are provided in Appendix E. $\square$

Finally, we make three remarks related to the results introduced in this section.

Remark 1. Both assumptions (i) and (ii) on the parameters in Theorem 3 will hold true when the number of communities is a constant, i.e., the size of each community is $C = \Theta(N)$ (as is the case when the population is well-mixed, or if we just consider a single community); assumption (i) does not require $C = \Theta(N)$. In our simulations, we observed empirically that assumption (ii) also holds when $C \ll N$; we do not have a formal proof of Theorem 3 for this case however.

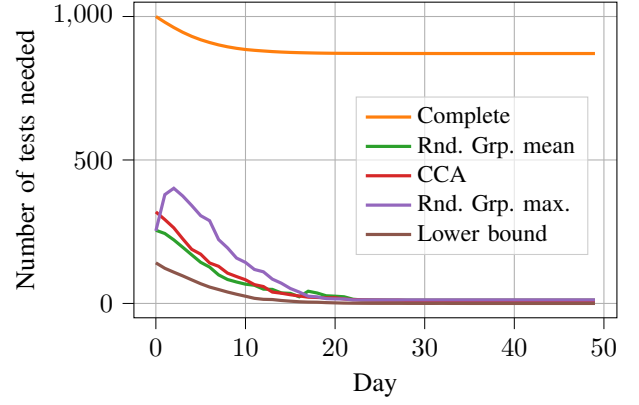
Remark 2. Our results hold not just for the specific model introduced in Section II (where in particular we assume symmetric intra and inter community transmissions) but for any underlying community structure where the two conditions (bounded prior vectors and the value of each prior not exceeding $1/2$) are satisfied. For example, one could have a model where an infected individual can transmit the infection only to a subset of his fellow community members with probability q_1 (he/she cannot transmit to the rest of the individuals in his/her community) and only to a subset of individuals outside his/her community with probability q_2 . For this example model, the conditions in Theorem 3 are sufficient to prove the two requirements on the prior vector.

Remark 3. Intervention is a crucial aspect for our results to hold true. Without intervention in our dynamic model, many of the prior probabilities would be greater than $1/2$ and our requirements on the prior vector would not be satisfied.

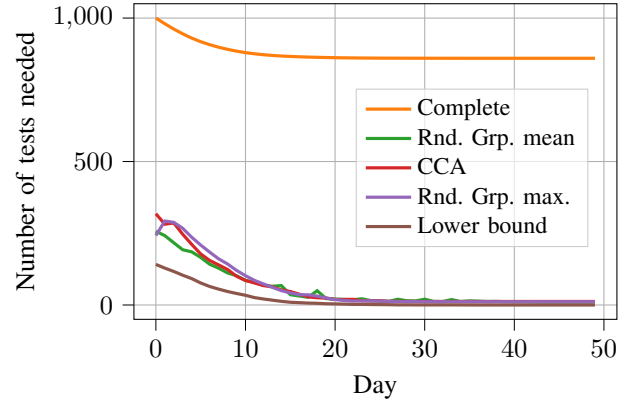
IV. NUMERICAL RESULTS

In this section, we show illustrative numerical results on the necessary and sufficient number of tests required for the discrete-time SIR stochastic block model. We next describe the experimental set-up.

- We simulate multiple instances (or *trajectories*) of the pipeline in Fig. 2 where the infections follow the discrete-time SIR stochastic block model $(N, C, p_{\text{init}}, q_1, q_2)$, and for different testing strategies. We simulate 200 trajectories and



(a) $(N, C, p_{\text{init}}, q_1, q_2) = (1000, 20, 0.02, 0.03, 0.0004)$.



(b) $(N, C, p_{\text{init}}, q_1, q_2) = (1000, 50, 0.02, 0.012, 0.0004)$.

Fig. 4: Experimental results. We plot the average number of tests required by each strategy to identify the infection statuses of all non-isolated individuals each day for 2 different sets of parameters.

in Fig. 4, plot the daily average of the quantities across these trajectories.

- For each of these testing strategies, we empirically find the number of tests required on each day to identify all the infections on the previous day. To do this, on each day for a given trajectory, we start with 1000 tests and decrease this number (at a certain granularity) until the testing strategy makes a mistake. The smallest number of tests for which the strategy worked is plotted.
- On the other hand, we also plot the *entropy lower bound* in Lemma 3; it is easy to estimate this for our model via Monte-Carlo approximations. This bound holds for any set of values for $p_i^{(t-1)}$, regardless of whether the conditions required for Theorem 3 hold or not. The reason we use the entropy bound instead of our lower bound in Theorem 2 is that the entropy bound was numerically observed to be larger. Indeed, the lower bound in Theorem 2 contains some accompanying hidden constants which are small when used for our particular choice of N .

We compare the following testing strategies in our numerical simulations.

- **Complete testing.** We test every non-isolated individual remaining in the population each day.

- **Coupon Collector Algorithm (CCA) from [14].** We showed the order-optimality of this algorithm for the dynamic testing problem at the beginning of Section III-C. In short, on each day, the CCA algorithm constructs a random non-adaptive test design which depends on $p_j^{(t)}$. The idea is to place objects which are less likely to be infected in more number of tests and vice-versa. We refer the reader to [14] for the exact description of the algorithm.
- **Random group testing for max probability (Rnd. Grp. max.)** Here we construct a randomized design assuming that each individual has a prior probability of infection $p_{\max}^{(t)}$. From Corollary 1, such a design must also work for the actual priors $p_j^{(t)}$. We construct a constant column-weight design (see e.g. [19]) where each individual is placed in $L = \lceil T \ln 2 / (N p_{\max}^{(t)}) \rceil$ tests. Such a test design achieves a vanishing probability of error with $O(N p_{\max}^{(t)} \log N)$ tests (see for example [19] for a proof), and hence is order-optimal under the conditions in Theorem 3.
- **Random group testing for mean probability (Rnd. Grp. mean)** Here we construct a randomized design assuming that each individual has a prior probability of infection $p_{\text{mean}}^{(t)}$, where $p_{\text{mean}}^{(t)}$ is defined as the mean prior probability of infection across all individuals. Unlike Rnd. Grp. max., there is no guarantee on how many tests are needed by such a design to identify the infection statuses of all individuals. However, the numerical results in Fig. 4 show that such a design requires fewer tests than CCA or Rnd. Grp. max. designs.

The numerical results in Fig. 4 are illustrated for two different parameter values of the discrete-time SIR stochastic block model. In both cases, we see that Rnd. Grp. mean requires the least number of tests to identify the infection statuses of all non-isolated individuals. Moreover, the number of tests required by all three testing strategies considered is much less than the number required by complete testing. In fact the numerics in Fig. 4 indicate that if we use a number of tests equal to $1/5$ of number of tests required for complete individual testing, all these algorithms would achieve the same performance as complete individual testing, at least for the particular examples that we considered.

A natural follow-up question to ask is if there is a systematic way to choose the number of tests that need to be administered, given the upper bounds discussed in Section III-B. In Appendix F, we discuss one such heuristic and show that it achieves close-to-complete-testing performance.

V. CONCLUSIONS AND OPEN QUESTIONS

In this work we proposed the problem of dynamic group testing which asks the question of how to continually test given that infections spread during the testing period. Our numerical results answer the question we started with – in the dynamic testing problem formulation, given a day of testing delay, is it possible to achieve close to complete testing performance with significantly fewer number of tests? The answer is yes, and in this paper we not only showed numerical evidence supporting this fact, but also gave theoretical bounds on the optimal number of tests needed in order to achieve this.

Alternatively, on a high-level, our work can also be interpreted as the first work that attempts to answer the following question: Given that we have some prior knowledge of how the infections are distributed today, and given information about how the infections spread over time, how many tests do I need tomorrow to identify all new infections? We provide upper and lower bounds on this number. Moreover, in Figure 4, we also estimate this number via Monte-Carlo simulations for a few example instances of the discrete-time block SIR model. These simulations can also be used to estimate other relevant quantities such as the maximum number of tests needed in a single day or the total number of tests needed until the epidemic dies out.

Although we gave theoretical bounds on the optimal number of tests needed each day, many open questions remain. In particular, it would be interesting to study the same problem when tests are noisy, or when we can use other models of group tests such as the one in [20], or simply when one cannot perfectly learn the states of all individuals on a testing day. In addition, it would also be of interest to study/use other test designs, for example like the ones in [19], [21]–[30]. Finally, it remains open to see how these results translate to the continuous-time SIR stochastic network model of [13].

VI. RELATED WORK

As stated in our introduction, our work shares similar goals with our prior work in [12], where we considered the well established continuous-time SIR stochastic network model (see [13]) and focused on how many tests to use and whom to test in order to track the infected individuals in the population. That work also explored how well one can learn the infected individuals given delayed test results, but gave no theoretical guarantees on the methods and did not consider intervention. Our discrete-time model for disease spread in this work, however, is more amenable to analysis and illustrates better the usefulness of group testing, being at the same time useful for practical reasons (more about this in Section II). We further note that our results are applicable to a more general set of SIR models as discussed in Section III-C, remark 2.

Our model is closely related to the independent cascade model (see for example [31] and references therein), studied in the context of influence maximization in social networks, where we can interpret influence/rumor propagation as infections in our context. A crucial difference of our model from this is that our model allows multiple opportunities of infections over time whereas the independent cascade model only allows one opportunity to "infect". Therefore, as is noted, in our model the infectious individuals remain infectious until recovered or isolated.

The work in [32] considers a discrete stochastic model for the progression of COVID-19 based on contact networks and leverages the model dynamics to inform a group test decoder; however their scope is different, as they test infrequently and thus infections are highly correlated, do not consider interventions, do not look for optimal group test designs, and do not provide theoretical guarantees on the number of tests needed.

Since we use the main principles of the SIR model our work is closely related to epidemic modeling. Works in epidemiology discuss the implications of testing and intervention for COVID-19 employing stochastic network models (see [8], [9] and references therein) but do not consider test designs that exploit the knowledge of the underlying dynamical system. Works in control theory (see [33] and references therein) consider deterministic SIR compartment models (at the population level) and focus on intervention schemes. Here we are interested in both testing and intervention and use an individual-level SIR model.

Our work can be positioned in the general context of community-aware group testing where infected are not independent, and correlations follow from the community structure. Our work in [18], [34] demonstrated that using a known community structure to design group testing strategies and decoding, can significantly extend the advantages of group testing by utilizing these structural dependencies. Concurrently, the works in [32], [35] proposed decoding algorithms that take the community structure into account. Following up on these works in the static case (without temporal dynamics), there have been other recent works with similar goals [36]–[38]. Our work also leverages a known community structure that informs the system dynamics as well as the group test designs.

Further related to static group testing is the work on graph-constrained group testing (see for example [39], [40]), which solves the problem of how to design group tests when there are constraints on which samples can be pooled together, provided in the form of a graph. In our case, no such constraints exist and individuals can be pooled together into tests freely.

VII. ACKNOWLEDGEMENTS

This work was supported in part by NSF grants #2007714, #2139304, #2146838 and #2146828. We also thank Katerina Argyraki for her ongoing support and the valuable discussions we have had about this project.

REFERENCES

- [1] C. Gollier and O. Gossner, “Group testing against covid-19,” April 2020, see <https://www.tse-fr.eu/articles/group-testing-against-covid-19>.
- [2] M. Broadfoot, “Coronavirus test shortages trigger a new strategy: Group screening,” May 2020, see <https://www.scientificamerican.com/article/coronavirus-test-shortages-trigger-a-new-strategy-group-screening2/>.
- [3] J. Ellenberg, “Five people. one test. this is how you get there.” *NYtimes*, May 2020.
- [4] C. Verdun *et al.*, “Group testing for sars-cov-2 allows up to 10-fold efficiency increase across realistic scenarios and testing strategies,” *medRxiv*, 2020.
- [5] S. Ghosh *et al.*, “Tapestry: A single-round smart pooling technique for covid-19 testing,” *medRxiv*, 2020.
- [6] 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,” *Annals of Internal Medicine*, vol. 173, pp. 262–267, Aug. 2020.
- [7] J. Taipale, I. Kontoyiannis, and S. Linnarsson, “Population-scale testing can suppress the spread of infectious disease,” 2021.
- [8] J. Taipale, P. Romer, and S. Linnarsson, “Population-scale testing can suppress the spread of covid-19,” *MedRxiv*, 2020.
- [9] T. Bergstrom, C. T. Bergstrom, and H. Li, “Frequency and accuracy of proactive testing for covid-19,” *medRxiv*, 2020.
- [10] M. Aldridge, O. Johnson, and J. Scarlett, “Group testing: an information theory perspective,” *CoRR*, vol. abs/1902.06002, 2019.
- [11] M. Aldridge and D. Ellis, “Pooled testing and its applications in the covid-19 pandemic,” *Pandemics: Insurance and Social Protection*, pp. 217–249, 2022.
- [12] S. R. Srinivasavaradhan, P. Nikolopoulos, C. Fragouli, and S. Diggavi, “An entropy reduction approach to continual testing,” in *2021 IEEE International Symposium on Information Theory (ISIT)*. IEEE, 2021, pp. 611–616.
- [13] I. Kiss, J. Miller, and P. Simon, *Mathematics of Epidemics on Networks*, 01 2017, vol. 46.
- [14] T. Li, C. L. Chan, W. Huang, T. Kaced, and S. Jaggi, “Group testing with prior statistics,” in *2014 IEEE International Symposium on Information Theory*, 2014, pp. 2346–2350.
- [15] R. Ben-Ami, A. Klochender, M. Seidel, T. Sido, O. Gurel-Gurevich, M. Yassour, E. Meshorer, G. Benedek, I. Fogel, E. Oiknine-Djian *et al.*, “Large-scale implementation of pooled rna extraction and rt-pcr for sars-cov-2 detection,” *Clinical Microbiology and Infection*, vol. 26, no. 9, pp. 1248–1253, 2020.
- [16] M. Chiani, G. Liva, and E. Paolini, “Identification-detection group testing protocols for covid-19 at high prevalence,” *Scientific Reports*, vol. 12, no. 1, p. 3250, 2022.
- [17] T. H. Cormen, C. E. Leiserson, R. L. Rivest, and C. Stein, *Introduction to algorithms*. MIT press, 2022.
- [18] P. Nikolopoulos, S. Rajan Srinivasavaradhan, T. Guo, C. Fragouli, and S. Diggavi, “Group testing for connected communities,” in *Proceedings of The 24th International Conference on Artificial Intelligence and Statistics*, vol. 130. PMLR, 2021, pp. 2341–2349. See also arXiv preprint arXiv:2007.08111.
- [19] O. Johnson, M. Aldridge, and J. Scarlett, “Performance of group testing algorithms with near-constant tests per item,” *IEEE Trans. Inf. Theory*, vol. 65, no. 2, pp. 707–723, 2019.
- [20] R. Gabrys, S. Pattabiraman, V. Rana, J. Ribeiro, M. Cheraghchi, V. Guruswami, and O. Milenkovic, “Ac-dc: Amplification curve diagnostics for covid-19 group testing,” 2020.
- [21] E. Price and J. Scarlett, “A fast binary splitting approach to non-adaptive group testing,” *arXiv preprint arXiv:2006.10268*, 2020.
- [22] S. Bondorf, B. Chen, J. Scarlett, H. Yu, and Y. Zhao, “Sublinear-time non-adaptive group testing with $\mathcal{O}(\log n)$ tests via bit-mixing coding,” *IEEE Transactions on Information Theory*, 2020.
- [23] M. Aldridge, L. Baldassini, and O. Johnson, “Group testing algorithms: Bounds and simulations,” *IEEE Transactions on Information Theory*, vol. 60, no. 6, pp. 3671–3687, 2014.
- [24] G. K. Atia and V. Saligrama, “Boolean compressed sensing and noisy group testing,” *IEEE Transactions on Information Theory*, vol. 58, no. 3, pp. 1880–1901, 2012.
- [25] J. Scarlett and V. Cevher, “Phase transitions in group testing,” in *Proceedings of the Twenty-Seventh Annual ACM-SIAM Symposium on Discrete Algorithms, SODA 2016, Arlington, VA, USA, January 10-12, 2016*. SIAM, 2016, pp. 40–53.
- [26] A. Coja-Oghlan, O. Gebhard, M. Hahn-Klimroth, and P. Loick, “Information-theoretic and algorithmic thresholds for group testing,” *IEEE Trans. Inf. Theory*, 2020.
- [27] 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, p. 3019–3035, 2014.
- [28] S. Cai, M. Jahangoshahi, M. Bakshi, and S. Jaggi, “Efficient algorithms for noisy group testing,” *IEEE Trans. Inf. Theory*, vol. 63, no. 4, p. 2113–2136, 2017.
- [29] H. A. Inan, P. Kairouz, M. Wootters, and A. Özgür, “On the optimality of the kautz-singleton construction in probabilistic group testing,” *CoRR*, vol. abs/1808.01457, 2018. [Online]. Available: <http://arxiv.org/abs/1808.01457>
- [30] K. Lee, R. Pedarsani, and K. Ramchandran, “Saffron: A fast, efficient, and robust framework for group testing based on sparse-graph codes,” in *2016 IEEE International Symposium on Information Theory (ISIT)*, 2016, pp. 2873–2877.
- [31] D. Kempe, J. Kleinberg, and É. Tardos, “Maximizing the spread of influence through a social network,” in *Proceedings of the ninth ACM SIGKDD international conference on Knowledge discovery and data mining*, 2003, pp. 137–146.
- [32] R. Goenka, S.-J. Cao, C.-W. Wong, A. Rajwade, and D. Baron, “Contact tracing enhances the efficiency of covid-19 group testing,” *arXiv preprint arXiv:2011.14186*, 2020.
- [33] T. G. Molnár, A. W. Singletary, G. Orosz, and A. D. Ames, “Safety-critical control of compartmental epidemiological models with measurement delays,” *IEEE Control Systems Letters*, vol. 5, no. 5, pp. 1537–1542, 2020.

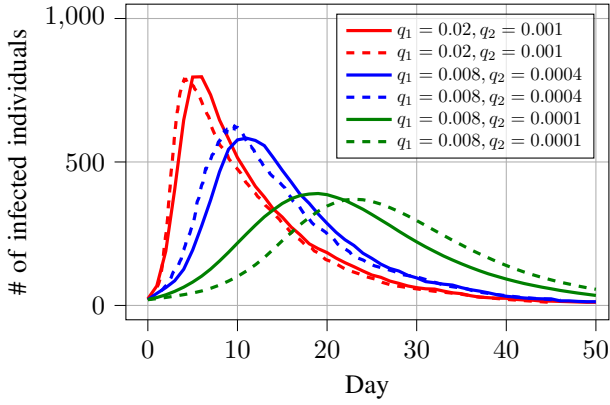


Fig. 5: Continuous vs discrete-time model. Continuous model in dashed and discrete model in solid curves. Recovery probability $r = 0.1$ in all cases.

- [34] P. Nikolopoulos, S. R. Srinivasavaradhan, T. Guo, C. Fragouli, and S. Diggavi, “Group testing for overlapping communities,” *arXiv preprint arXiv:2012.02804*, IEEE ICC, 2021.
- [35] J. Zhu, K. Rivera, and D. Baron, “Noisy pooled pcr for virus testing,” *arXiv preprint arXiv:2004.02689*, 2020.
- [36] S. Ahn, W.-N. Chen, and A. Ozgur, “Adaptive group testing on networks with community structure,” *arXiv preprint arXiv:2101.02405*.
- [37] B. Arasli and S. Ulukus, “Group testing with a graph infection spread model,” *arXiv preprint arXiv:2101.05792*, 2021.
- [38] P. Bertolotti and A. Jadbabaie, “Network group testing,” *arXiv preprint arXiv:2012.02847*, 2020.
- [39] M. Cheraghchi, A. Karbasi, S. Mohajer, and V. Saligrama, “Graph-constrained group testing,” *IEEE Transactions on Information Theory*, vol. 58, no. 1, pp. 248–262, 2012.
- [40] A. Karbasi and M. Zadimoghaddam, “Sequential group testing with graph constraints,” in *2012 IEEE information theory workshop*. Ieee, 2012, pp. 292–296.
- [41] W. H. Bay, E. Price, and J. Scarlett, “Optimal non-adaptive probabilistic group testing requires $\theta(\min\{k \log n, n\})$ tests,” *arXiv e-prints*, 2020.
- [42] A. Coja-Oghlan, O. Gebhard, M. Hahn-Klimroth, and P. Loick, “Optimal group testing,” ser. *Proceedings of Machine Learning Research*, J. Abernethy and S. Agarwal, Eds., vol. 125, Jul. 2020, pp. 1374–1388.
- [43] M. Aldridge, “Individual testing is optimal for nonadaptive group testing in the linear regime,” *IEEE Trans. Inf. Theory*, vol. 65, no. 4, 2019.

APPENDIX A

COMPARISON OF DISCRETE AND CONTINUOUS-TIME SIR MODELS

The well-studied continuous-time SIR stochastic network model from [13] has been the main motivation for our discrete-time SIR stochastic block model. In fact, the discrete-time SIR stochastic block model described in Section II-B can be considered as a discretized version of the continuous-time SIR stochastic network model over the weighted graph, where 2 individuals belonging to the same community are connected by an edge with weight q_1 and 2 individuals belonging to different communities are connected by an edge with weight q_2 , and recoveries occur at the rate r/day – i.e., an infected individual transmits the disease to a susceptible individual in the same community at the rate q_1/day and to a susceptible individual in a different community at the rate q_2/day . In Fig. 5, we compare the continuous-time model above and the discrete-time model for a few example values of q_1 , q_2 and r for illustration.

We make a few observations:

- The progression of the disease in the discrete-time and

continuous-time models, though not identical, follow a similar pattern, justifying the use of the discrete-time model.

- In both the models, $1/q_1$ is the expected time for an infected individual to transmit the disease to a susceptible individual in the same community, $1/q_2$ is the expected time for an infected individual to transmit the disease to a susceptible individual in a different community and $1/r$ is the expected time for an infected individual to recover.

- In the continuous-time model, an individual can get infected and recovered in the same day, whereas this is not possible in our discrete-time model (infected individuals can recover starting from the day after they are infected).

APPENDIX B

PRELIMINARY: REVIEW OF RESULTS FROM STATIC GROUP TESTING

Traditional group testing typically assumes a population of N individuals out of which some are infected. Three infection models are typically considered: (i) in the *combinatorial priors model*, a fixed number of infected individuals k , are selected uniformly at random among all sets of size k ; (ii) in the *i.i.d probabilistic priors model*, each individual is i.i.d infected with probability p ; (iii) in the *non-identical probabilistic priors model*, each item i is infected independently of all others with prior probability p_i , so that the expected number of infected members is $\bar{k} = \sum_{i=1}^N p_i$ [14]. In this paper we mostly use results that apply to the last case.

A group test 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 testing output y_τ takes a binary value calculated as $y_\tau = \bigvee_{i \in \mathcal{D}_\tau} U_i$ ³, where $\bigvee$ stands for the OR operator (disjunction) and $\mathcal{D}_\tau$ is the group of people participating in the test.

The usual goal in static group testing is to design a testing algorithm that is able to identify all infection statuses $\mathbf{U} = (U_1, \dots, U_N)$. These algorithms can be adaptive or non-adaptive. 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. Non-adaptive 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 non-zero elements determine the set $\mathcal{D}_\tau$. Although adaptive testing uses less tests than non-adaptive, non-adaptive testing is often more practical as all tests can be executed in parallel.

The main challenge in static group testing is the number of group tests $T = T(N)$ needed to identify the infected members without error or with high probability. In the following, we present some well established results that we reuse in our work (hereafter we use typical asymptotic notations; i.e., $f(n) = O(g(n))$, $f(n) = \Omega(g(n))$, or $f(n) = \Theta(g(n))$ respectively means that a $f(n) \leq cg(n)$, $f(n) \geq cg(n)$,

³We assume that the tests are noiseless here, for simplicity. The group testing literature also extensively studies the case when the testing output is noisy.

$c_1 g(n) \leq f(n) \leq c_2 g(n)$ asymptotically, where c, c_1, c_2 are universal constants):

- For the probabilistic model (ii), any non-adaptive algorithm with a success probability bounded away from zero as $N \rightarrow \infty$ must have $T = \Omega(\min\{\bar{k} \log N, N\})$ [41, Theorem 1], [42]. This means that either any non-adaptive group testing with a number of tests $O(\bar{k} \log N)$ is order optimal, or individual testing is order optimal⁴. In particular, random test designs, such as i.i.d. Bernoulli [23]–[25] and near-constant tests-per-item [19], [26] have been proved to be order-optimal in a sparse regime where $\bar{k} = \Theta(N^\alpha)$ and $\alpha \in (0, 1)$. In fact, in the same regime, [42] has provided the precise constants for optimal non-adaptive group testing. Conversely, classic individual testing has been proved to be optimal in the linear ($\bar{k} = \Theta(N)$) [43] and the mildly sublinear regime ($\bar{k} = \omega(\frac{N}{\log N})$) [41].

- For the probabilistic model (iii), a lower bound for the number of tests needed is given by the entropy, stated below:

Lemma 3 (Entropy lower bound). *Consider the non-identical probabilistic priors model of static group testing, where each individual $i \in [N]$ is infected independently with probability p_i . The number of tests T needed by a non-adaptive algorithm to identify the infection status of all individuals with a vanishing probability of error satisfies*

$$T \geq \sum_{i=1}^N h_2(p_i),$$

where $h_2(\cdot)$ is the binary entropy function.

See Appendix A in [14] for a proof. On the algorithmic side, two 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 $4e(1 + \delta)\bar{k} \ln N$ tests to achieve an error probability no larger than $2N^{-\delta}$ whenever $p_i \leq 1/2$ [14], [27].

Distinction from traditional group testing. In this paper, we focus on probabilistic infections and non-adaptive test designs, but we differ from traditional group testing in two ways:

(a) Traditional group testing examines a static scenario, where the state of individuals is fixed (infected or not); we are instead interested in a dynamic scenario, where the state of an individual may change, even during the test period. This is particularly true since test results may not be available instantaneously but instead with a delay (e.g. after one day).
 (b) The infection probabilities p_i are not independent; instead, they are correlated where the correlation is induced by the underlying community structure and dynamic infection spread model we consider (for instance, two individuals who live in the same household are more likely to be both infected or not). This implies that U_i and U_j are not independent, as is the assumption in traditional group testing.

⁴The achievability and converse results provided here are usually proved for combinatorial model (i) (a summary can be found in [21]), but they are directly applicable to model (ii) by considering $p = k/N$ (see Theorem 1.7 and Theorem 1.8 in [10] or [41]).

APPENDIX C PROOF OF LEMMA 1

The optimality of the MAP decoder is a standard result in statistics and signal processing. We however give the proof in the context of our problem, for completeness.

Lemma 1 (Optimality of MAP decoder). *For given test matrix G and priors $\mathbf{p}$, the corresponding MAP decoder minimizes the probability of error for the test matrix under the given priors, i.e.,*

$$\mathbb{P}_{err}(G, R_{map}(\cdot; G, \mathbf{p}); \mathbf{p}) \leq \mathbb{P}_{err}(G, R; \mathbf{p}) \quad \forall R.$$

Proof. As stated at the beginning of Section III, the Probability of error for a test matrix, decoder pair under given the priors is

$$\begin{aligned} \mathbb{P}_{err}(G, R; \mathbf{p}) &\triangleq \mathbb{E}_{\mathbf{U} \sim \mathbf{p}} \mathbb{1}\{R(G(\mathbf{U})) \neq \mathbf{U}\} \\ &= \mathbb{E}_{\mathbf{Y}} \mathbb{E}_{\mathbf{U}|\mathbf{Y}} \mathbb{1}\{R(G(\mathbf{U})) \neq \mathbf{U}\}, \end{aligned}$$

where $\mathbf{Y}$ is the set of test results. For the MAP decoder, the term inside $\mathbb{E}_{\mathbf{Y}}$ is

$$\begin{aligned} &\mathbb{E}_{\mathbf{U}|\mathbf{Y}} \mathbb{1}\{R_{map}(\mathbf{Y}; G, \mathbf{p}) \neq \mathbf{U}\} \\ &= \mathbb{E}_{\mathbf{U}|\mathbf{Y}} \mathbb{1}\left\{ \arg \max_{\mathbf{U}: G(\mathbf{U})=\mathbf{Y}} \mathbb{P}(\mathbf{U}; \mathbf{p}) \neq \mathbf{U} \right\} \\ &= \sum_{\mathcal{D} \subseteq [N]} \mathbb{P}(\mathbf{U}(\mathcal{D})|\mathbf{Y}; \mathbf{p}) \\ &\quad \cdot \mathbb{1}\left\{ \arg \max_{\mathbf{U}(\mathcal{D}): G(\mathbf{U}(\mathcal{D}))=\mathbf{Y}} \mathbb{P}(\mathbf{U}(\mathcal{D}); \mathbf{p}) \neq \mathbf{U}(\mathcal{D}) \right\} \\ &= 1 - \mathbb{P}\left(\arg \max_{\mathbf{U}: G(\mathbf{U})=\mathbf{Y}} \mathbb{P}(\mathbf{U}|\mathbf{Y}; \mathbf{p}) \right) \\ &= 1 - \max_{\mathbf{U}: G(\mathbf{U})=\mathbf{Y}} \frac{\mathbb{P}(\mathbf{U}; \mathbf{p})}{\mathbb{P}(\mathbf{Y})}. \end{aligned} \tag{5}$$

Similarly, for any decoder R , we have

$$\begin{aligned} &\mathbb{E}_{\mathbf{U}|\mathbf{Y}} \mathbb{1}\{R(\mathbf{Y}) \neq \mathbf{U}\} \\ &= \sum_{\mathcal{D} \subseteq [N]} \mathbb{P}(\mathbf{U}(\mathcal{D})|\mathbf{Y}; \mathbf{p}) \cdot \mathbb{1}\{R(\mathbf{Y}) \neq \mathbf{U}\} \\ &= 1 - \mathbb{P}(R(\mathbf{Y})|\mathbf{Y}; \mathbf{p}) \geq 1 - \max_{\mathbf{U}: G(\mathbf{U})=\mathbf{Y}} \frac{\mathbb{P}(\mathbf{U}; \mathbf{p})}{\mathbb{P}(\mathbf{Y})}. \end{aligned} \tag{6}$$

Comparing (5) and (6) concludes the proof. $\square$

APPENDIX D PROOF OF LEMMA 2

Lemma 2. *Consider a test matrix G and priors $\mathbf{p}$, and let $p_i \leq 0.5$ for all i . Suppose the corresponding MAP decoder is erroneous when identifying the defective set $\mathcal{D}$. Then the MAP decoder is also erroneous when identifying the defectives set $\mathcal{D} \cup \{j\}$, i.e.,*

$$\begin{aligned} &\mathbb{1}\{R_{map}(G(\mathbf{U}(\mathcal{D} \cup \{j\})); G, \mathbf{p}) \neq \mathbf{U}(\mathcal{D} \cup \{j\})\} \\ &\geq \mathbb{1}\{R_{map}(G(\mathbf{U}(\mathcal{D})); G, \mathbf{p}) \neq \mathbf{U}(\mathcal{D})\}. \end{aligned}$$

Proof. We first state the trivial case where $\mathbb{1}\{R_{map}(G(\mathbf{U}(\mathcal{D})); G, \mathbf{p}) \neq \mathbf{U}(\mathcal{D})\} = 0$. Under that assumption, the inequality of Lemma 2 always holds.

We then consider the case where $\mathbb{1}\{R_{\text{map}}(G(\mathbf{U}(\mathcal{D})); G, \mathbf{p}) \neq \mathbf{U}(\mathcal{D})\} = 1$, i.e., the MAP decoder makes an error when the defective set is $\mathcal{D}$. In that case, one of the two situations is possible:

- (1) there exists some set $\mathcal{D}' \neq \mathcal{D}$, such that $G(\mathbf{U}(\mathcal{D}')) = G(\mathbf{U}(\mathcal{D}))$ and $\mathbb{P}(\mathbf{U}(\mathcal{D}'); \mathbf{p}) > \mathbb{P}(\mathbf{U}(\mathcal{D}); \mathbf{p})$ or
- (2) there exists some set $\mathcal{D}' \neq \mathcal{D}$, such that $G(\mathbf{U}(\mathcal{D}')) = G(\mathbf{U}(\mathcal{D}))$ and $\mathbb{P}(\mathbf{U}(\mathcal{D}'); \mathbf{p}) = \mathbb{P}(\mathbf{U}(\mathcal{D}); \mathbf{p})$ and $\mathcal{D}'$ is lexicographically earlier than $\mathcal{D}$.

Hence MAP identifies incorrectly $\mathcal{D}'$ as the defective set given that $\mathcal{D}$ was the true defective set. We prove assuming that the first situation occurred; the proof follows identical arguments for the second situation.

Now, we consider two different cases for individual j that is added to $\mathcal{D} \cup \{j\}$:

- (i) If $j \notin \mathcal{D}'$, then from our assumption in (1), notice that the defective set $\mathcal{D}' \cup \{j\}$ explains the test results of $\mathcal{D} \cup \{j\}$ – $\mathcal{D}'$ gives the same test results as $\mathcal{D}$ and the extra individual j added to both the sets will still give the same results. We next claim that $\mathbb{P}(\mathbf{U}(\mathcal{D}' \cup \{j\}); \mathbf{p}) > \mathbb{P}(\mathbf{U}(\mathcal{D} \cup \{j\}); \mathbf{p})$, and consequently the MAP decoder will fail to correctly identify the defective set $\mathcal{D}' \cup \{j\}$. Now to prove our claim, we start with our assumption (1), i.e.,

$$\begin{aligned}
& \mathbb{P}(\mathbf{U}(\mathcal{D}'); \mathbf{p}) > \mathbb{P}(\mathbf{U}(\mathcal{D}); \mathbf{p}) \\
& \implies \prod_{i \in \mathcal{D}'} p_i \prod_{l \in [N] \setminus \mathcal{D}'} (1 - p_l) > \prod_{i \in \mathcal{D}} p_i \prod_{l \in [N] \setminus \mathcal{D}} (1 - p_l) \\
& \stackrel{(a)}{\implies} (1 - p_j) \prod_{i \in \mathcal{D}'} p_i \prod_{l \in [N] \setminus \mathcal{D}' \cup \{j\}} (1 - p_l) \\
& > (1 - p_j) \prod_{i \in \mathcal{D}} p_i \prod_{l \in [N] \setminus \mathcal{D} \cup \{j\}} (1 - p_l) \\
& \stackrel{(b)}{\implies} p_j \prod_{i \in \mathcal{D}'} p_i \prod_{l \in [N] \setminus \mathcal{D}' \cup \{j\}} (1 - p_l) \\
& > p_j \prod_{i \in \mathcal{D}} p_i \prod_{l \in [N] \setminus \mathcal{D} \cup \{j\}} (1 - p_l) \\
& \stackrel{(c)}{\implies} \prod_{i \in \mathcal{D}' \cup \{j\}} p_i \prod_{l \in [N] \setminus \mathcal{D}' \cup \{j\}} (1 - p_l) \\
& > \prod_{i \in \mathcal{D} \cup \{j\}} p_i \prod_{l \in [N] \setminus \mathcal{D} \cup \{j\}} (1 - p_l) \\
& \implies \mathbb{P}(\mathbf{U}(\mathcal{D}' \cup \{j\}); \mathbf{p}) > \mathbb{P}(\mathbf{U}(\mathcal{D} \cup \{j\}); \mathbf{p}),
\end{aligned}$$

where in (a) we take out the term corresponding to j , also we use the fact that $j \notin \mathcal{D}$ and $j \notin \mathcal{D}'$; (b) follows from multiplying both sides with $p_j/1 - p_j$; in (c) we push the p_j term into the first product term.

- (ii) If $j \in \mathcal{D}'$, we again first note that the defective set $\mathcal{D}' \cup \{j\} = \mathcal{D}'$ explains the test results of $\mathcal{D} \cup \{j\}$. We next claim that $\mathbb{P}(\mathbf{U}(\mathcal{D}'); \mathbf{p}) > \mathbb{P}(\mathbf{U}(\mathcal{D} \cup \{j\}); \mathbf{p})$, and consequently the MAP decoder will fail to correctly identify the defective set $\mathcal{D}' \cup \{j\}$. Now to prove our claim, we start with our assumption (1), i.e.,

$$\begin{aligned}
& \mathbb{P}(\mathbf{U}(\mathcal{D}'); \mathbf{p}) > \mathbb{P}(\mathbf{U}(\mathcal{D}); \mathbf{p}) \\
& \implies \prod_{i \in \mathcal{D}'} p_i \prod_{l \in [N] \setminus \mathcal{D}'} (1 - p_l) > \prod_{i \in \mathcal{D}} p_i \prod_{l \in [N] \setminus \mathcal{D}} (1 - p_l)
\end{aligned}$$

$$\begin{aligned}
& \stackrel{(a)}{\implies} p_j \prod_{i \in \mathcal{D}' \setminus \{j\}} p_i \prod_{l \in [N] \setminus \mathcal{D}'} (1 - p_l) \\
& > (1 - p_j) \prod_{i \in \mathcal{D}} p_i \prod_{l \in [N] \setminus \mathcal{D} \cup \{j\}} (1 - p_l) \\
& \stackrel{(b)}{\implies} p_j \prod_{i \in \mathcal{D}' \setminus \{j\}} p_i \prod_{l \in [N] \setminus \mathcal{D}'} (1 - p_l) \\
& > p_j \prod_{i \in \mathcal{D}} p_i \prod_{l \in [N] \setminus \mathcal{D} \cup \{j\}} (1 - p_l) \\
& \stackrel{(c)}{\implies} \prod_{i \in \mathcal{D}'} p_i \prod_{l \in [N] \setminus \mathcal{D}'} (1 - p_l) \\
& > \prod_{i \in \mathcal{D} \cup \{j\}} p_i \prod_{l \in [N] \setminus \mathcal{D} \cup \{j\}} (1 - p_l) \\
& \implies \mathbb{P}(\mathbf{U}(\mathcal{D}'); \mathbf{p}) > \mathbb{P}(\mathbf{U}(\mathcal{D} \cup \{j\}); \mathbf{p}),
\end{aligned}$$

where in (a) we take out the term corresponding to j , also note that $j \notin \mathcal{D}$ but $j \in \mathcal{D}'$; (b) follows from the fact that $1 - p_j \geq p_j$ when $p_j \leq 0.5$, so we can replace the $(1 - p_j)$ term on the right-hand side by p_j without affecting the inequality; in (c) we push the p_j term into the first product term. $\square$

APPENDIX E

AUXILIARY RESULTS FOR THEOREM 3

In this section we prove some auxiliary statements about functions $f_1(x)$, $f_2(x)$ and $f_3(x)$ that are used at the end of the proof of Theorem 3:

- $f_1(x) = \frac{1 - \kappa x}{1 - x}$ is increasing for $\kappa \in (0, 1)$, because $f'_1(x) = -\frac{\kappa - 1}{(x - 1)^2} > 0$.
- $f_2(x) = (1 - q_2)^x$ is decreasing for $q_2 \in (0, 1)$, because $f'_2(x) = \ln(1 - q_2)(1 - q_2)^x < 0$.
- $f_3(x) = \frac{1 - c_1^x}{1 - c_2^x}$ is decreasing for $q_1 \geq q_2$, because of the following: Let $c_1 = 1 - q_1$ and $c_2 = 1 - q_2$, so that $c_2 \geq c_1$. Then,

$$\begin{aligned}
f'_3(x) &= \frac{1}{(1 - c_2^x)^2} ((1 - c_1)^x c_2^x \ln c_2 - (1 - c_2^x) c_1^x \ln c_1) \\
&= \frac{1}{x(1 - c_2^x)^2} ((1 - c_1)^x c_2^x \ln c_2^x - (1 - c_2^x) c_1^x \ln c_1^x) \\
&= \frac{(1 - c_1)^x (1 - c_2)^x}{x(1 - c_2^x)^2} \left(\frac{c_2^x \ln c_2^x}{(1 - c_2)^x} - \frac{c_1^x \ln c_1^x}{(1 - c_1)^x} \right) \stackrel{(a)}{\leq} 0,
\end{aligned}$$

where (a) follows from the fact that $c_2 \leq c_1$ and the function $g(c) = \frac{c \ln c}{1 - c}$ is non-increasing for $c \in (0, 1)$. The latter can be seen by taking the derivative $g'(c) = \frac{\ln c - c + 1}{(1 - c)^2}$, which is always non-positive for $c \in (0, 1)$, as $\ln c \leq c - 1$.

APPENDIX F

A HEURISTIC FOR DYNAMIC GROUP TESTING

Given the results and discussion in Section IV, a natural question to ask is if one could use a number of tests based on the upper bounds discussed in Section III-B. In particular, we focus on the upper bound for CCA which implies that CCA achieves a probability of error less than $2N^{-\delta}$ with a number of tests at most $4e(1 + \delta)Np_{\text{mean}} \log N$ (see Theorem 3 in [14]). Note that the probability of error is small, but not zero, for finite values of N . Here, we use a number of

tests equal to $12eNp_{\text{mean}} \log N$ each day (corresponding to an error probability less than $2N^{-2}$) and plot the number of errors made by each of the three test designs considered in Section IV. The experimental set-up is as follows:

(i) We maintain an estimate of the probability $p_j^{(t)}$ that a susceptible individual belonging to community j becomes infected on day t (see Section II-B for the precise definition), for each community j , and for each day t .

(ii) At the beginning of day t , we obtain the results of the tests administered on day $t-1$. From these results, we form an estimate $\hat{U}_i^{(t-1)}$ of the infection statuses $U_i^{(t-1)}$ of individual i at the beginning of day $t-1$, for each i . In order to learn the statuses, we use the *Definite Defective* (DD) decoder (see Section 2.4 in [10]) which is guaranteed to have no false positives. Indeed, one could use more sophisticated decoders, such as ones based on loopy belief propagation. However, these decoders potentially give rise to both false positives and false negatives, resulting in an unfair comparison across different algorithms⁵.

(iii) We isolate all individuals i where $\hat{U}_i^{(t-1)} = 1$.

(iv) We update $p_j^{(t-1)}$ using our estimates $\hat{U}_i^{(t-1)}$.

(v) Using our estimates of $p_j^{(t-1)}$, we estimate the value of $p_{\text{mean}}^{(t)}$ and choose a number of tests

$$T = \min\{12eN^{(t)}p_{\text{mean}}^{(t)} \log N^{(t)}, N^{(t)}\},$$

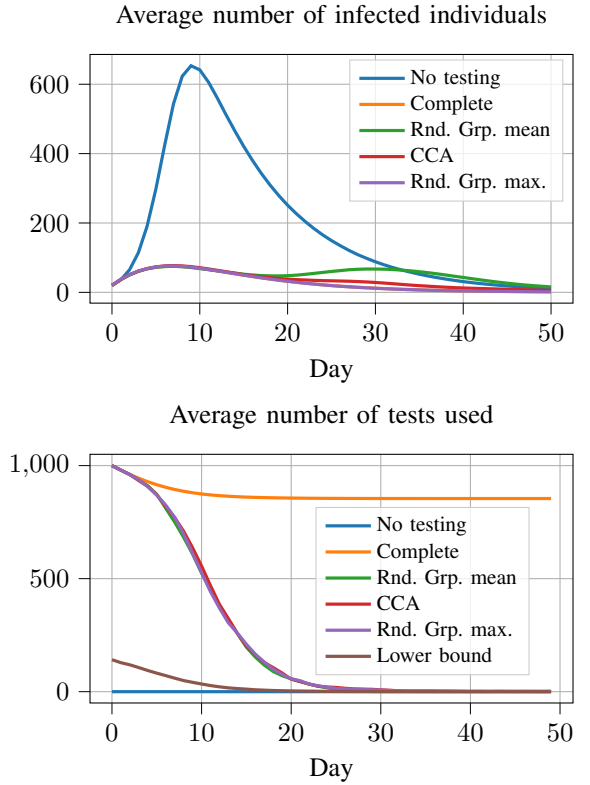
where $N^{(t)}$ is the current number of non-isolated individuals in the community. We next construct a testing matrix with T tests and administer these tests. For complete testing, we use $T = N^{(t)}$.

(vi) Steps (ii) – (v) repeat each day.

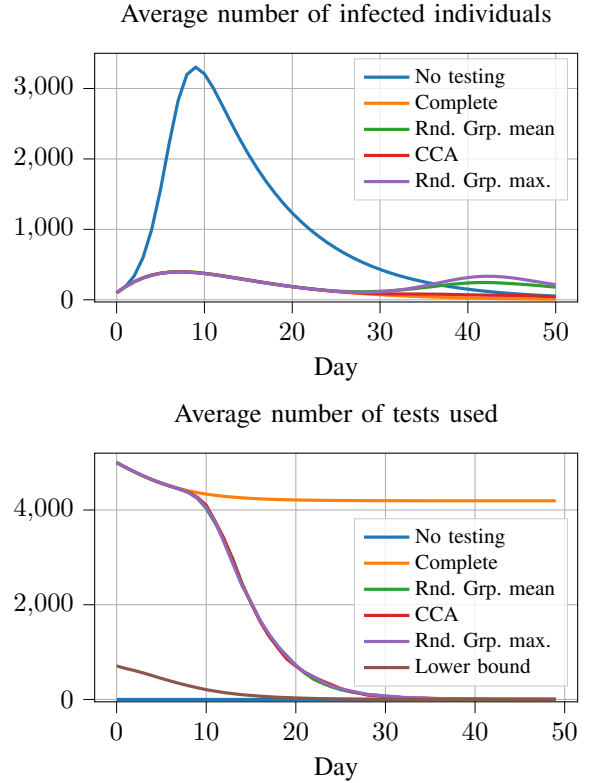
Given the above set-up, Figure 6 compares the performance of the test designs described in Section IV. We make a few observations:

- We see that the algorithms do not always attain the performance of complete testing. This is due to the fact that for finite N , the probability of error is non-zero. However, as seen from Figure 6, the performance of CCA improves as N increases. On the other hand, Rnd. Grp. max. has the opposite trend; this is not surprising since the number of tests was chosen based on the upper bound for CCA and as a result there is no guarantee that the same number of tests is sufficient for Rnd. Grp. max.
- In comparison to the plots in fig. 4, the number of tests used here is much higher during the initial few days, which indicates the looseness of the upper bound; it remains open to show tighter upper bounds for these algorithms.
- Suppose we make an error when identifying the infection status on a particular day t , the estimates of $p_j^{(t-1)}$ are not exact, which in turn leads to potentially insufficient choices for the number of tests needed for subsequent days and inaccurate test designs. This drives an error accumulation and as a result the later days are more prone to error, as also seen in Figure 6.

⁵Indeed, this begs the very complicated comparison between the impact of false positives and false negatives, which we avoid for the sake of simplicity.



(a) $(N, C, p_{\text{init}}, q_1, q_2) = (1000, 50, 0.02, 0.012, 0.0004)$.



(b) $(N, C, p_{\text{init}}, q_1, q_2) = (5000, 50, 0.02, 0.012, 8 \times 10^{-5})$.

Fig. 6: Experimental results for the heuristic procedure described in Appendix F. We plot the average number of infected individuals and the number of tests used as a function of time (in days). For comparison, we also plot the performance when no one is tested (no testing) and when everyone is tested (Complete).