

**Managing the Risk From COVID-19
During a Return to On-Site University Research**

Contact: Gordon Long - glong@mitre.org

JSR-20-NS1

July 2, 2020

Distribution Statement A: Approved for Public Release. Distribution is unlimited.

JASON
The MITRE Corporation
7515 Colshire Drive
McLean, Virginia 22102
(703) 983-6997

Executive Summary

JASON charged itself to assess risks and best practices for restarting university research programs. Three elements comprise the charge:

1. Understand the primary sources of risk and how they interact with the university environment;
2. Suggest modifications to existing continuity-of-operations plans;
3. Identify the relevant information for collection from the personnel involved.

JASON members began meeting weekly on April 10. Three subgroups (Testing, Modeling, and Operations and Safety) met separately to develop results for weekly meetings. We have consulted the literature, much of it developed in just the last weeks, and relied on our expertise as university faculty to guide choices. Many members were involved in planning research restarts at their own universities.

The report is organized around eight questions that an administrator for a research institution may have:

1. What are the relevant characteristics of COVID-19?
2. How is the risk of airborne transmission?
3. What is the role of diagnostic testing?
4. How can health screening reduce spread?
5. How does one prevent a super-spreader event?
6. Is the campus an island from the community?
7. What operational policies are recommended?
8. How can institutions make risk-informed decisions going forward?

Each section of the report seeks to answer these questions through a narrative and concludes with findings and recommendations.

Key findings and recommendations

After months of allowing only essential personnel on campus, universities have begun to ramp up research and allow typically a fraction of their researchers back on campus. The operations must evolve to keep the basic reproduction number, R , less than one in order to allow future increases in personnel density. This report outlines several different means of lowering the basic reproduction number, which are summarized in this section.

A low basic reproduction number relies on individuals recognizing the danger of COVID-19 and taking the recommended actions: washing hands, wearing face masks, keeping a minimum six foot separation from others, minimizing multiple occupancy in rooms, and tracking their own health. Universities must encourage these actions through signage, training, and modeling of good behavior by leaders. Face masks are particularly important.

Key Finding: Mask use can be highly effective as one component of risk reduction strategies for COVID-19 infection and transmission. However, mask effectiveness is variable, depending on the materials, designs and user discipline in wearing the masks.

Key Recommendation: Universities should provide masks that meet demonstrated technical performance at the levels needed, even if the level of performance exceeds that required by the city or state. Training should be provided on how to properly wear masks.

Key Finding: The use of a campus-wide “infographic” or “dashboard” showing the on-campus population, virus testing statistics, and information on the compliance with COVID-19 rules will create a shared situational awareness.

In this pandemic, universities are not islands and the reproduction number on-campus will not be very different than that of the surrounding community. If the R of the surrounding community changes, the university may need to change its on-campus density.

Key Finding: Universities will influence and be subject to disease dynamics of the larger communities within which they are embedded.

Key Recommendation: Universities should engage with state and local officials to understand the exposure of their personnel both on and off cam-

pus.

Universities can make use of daily symptom attestation to detect emerging cases and track the health of the population. Daily symptom attestation can reduce disease spread beyond what can be accomplished with individual mitigation behaviors such as wearing masks.

Key Recommendations: Develop a procedure for daily symptom attestation. A cellphone app for attestation offers the valuable opportunity to automatically determine if a probability-of-illness threshold has been reached. If symptoms exceed an established risk threshold, the individual should not come to work until they are confirmed to be clear of disease (e.g., negative test, quarantine period). Assume a high false-negative rate for symptom attestation during planning. Develop a procedure for rapid contact tracing when a researcher tests positive for COVID-19.

Unless a substantial fraction of the population can be tested on a daily basis and with fast turnaround times, virus testing will not help much in preventing transmission on campus. Testing can also serve other purposes, such as surveillance, but each purpose has particular requirements that must be addressed in advance.

Key Recommendation: When planning a virus testing program, make sure to understand the implications of false positives and false negatives to ensure the testing program yields meaningful results.

Super-spreaders, individuals with viral loads up to 10,000 times higher than average, cause infection outbreaks by rapidly infecting a large number of people.

Key Finding: A tracking system should respond quickly enough that once a symptomatic individual is detected, all contacts with that individual for the past 3-5 days are identified, notified and isolated in less than a day. Without this response speed, it may be difficult to stay ahead of the spread based on symptoms alone.

Key Recommendation: Limit the number of persons that an individual can come into contact with (e.g., through room and floor occupancy limits) to cap the size of super-spreading events.

Key Recommendation: When inexpensive, rapid virus tests become com-

mercially available, use them to test daily at the start of the work day to detect pre-symptomatic individuals, especially those with high viral titers.

Aerosols are an important means of transmission in laboratories and other enclosed spaces. In addition to wearing masks, minimizing double occupancy, and maintaining distance, buildings' HVAC systems can play a role in mitigating transmission.

Key Recommendation: Laboratory directors should consult their university's facilities and health and safety group on airflow in their labs to ensure at least 4 air changes per hour (ACH) is taking place and to increase the flow rate if the lab has more than single occupancy.

Contents

1	Introduction	9
2	What are the relevant characteristics of COVID-19?	12
3	What is the risk of airborne transmission?	15
3.1	Risk reduction	16
3.1.1	Physical distancing	17
3.1.2	Masks and eye protection	18
3.1.3	Reducing the occupancy of the room	18
3.1.4	Modifications to HVAC systems	18
3.1.5	Wearing masks as source control	21
3.1.6	Reducing droplet generating activities	21
4	What is the role of diagnostic testing?	24
4.1	The kinds of tests	24
4.1.1	Viral RNA testing	24
4.1.2	Viral antigen testing	26
4.1.3	Antibody testing	27
4.2	The functions of testing	28
4.2.1	Diagnosis	28

4.2.2	Screening	28
4.2.3	Monitoring	31
5	How can health screening reduce spread?	34
6	How does one prevent a super-spreader event?	38
7	Is the campus an island from the community?	43
7.1	Island model	43
7.2	Archipelago model	44
8	What operational policies are recommended?	46
8.1	Communication	46
8.2	Basic source of risk	46
8.3	Organizational methods to reduce dose	47
8.4	Respiratory masks	47
8.4.1	Mask technical standards	48
8.5	Eye protection	51
8.6	Physical distancing and its limits	51
8.7	Compliance	52
9	How can institutions make risk-informed decisions?	54

A	Lessons from the 1918 pandemic	56
B	Strategies to reduce R	59
B.1	Increasing the rate of removal	59
B.1.1	Screening via symptom attestation	60
B.1.2	Testing	61
B.1.3	Pulsed testing	63
B.2	Sub-population	64
C	Shift work and community interactions	66
C.1	Spread of a virus	66
C.2	Basic Model: SEIR for a Single Group	67
C.2.1	Time Course for Development of Symptoms of COVID-19	68
C.3	Growth of an Epidemic and Lockdown	69
C.3.1	An earlier lockdown	71
C.4	Dynamics following a return to work	71
C.5	Model of Karin et al.: A “4:10” Strategy,	73
C.5.1	The idea of the “4-10” work cycle	73
C.5.2	Two halves of the community in “4-10” cycles	75
D	Daily testing at the start of the workday	76
E	Pooled testing	78

E.1	Incorporating test imperfections	80
F	Marginal risk from COVID-19	83
G	Expected utility for decision-making	87
H	Simple analysis of the impact of false-positive tests	94

1 Introduction

Universities and national laboratories are ramping-up research operations following several months of curtailed activity during the early stages of the COVID-19 pandemic. In addition Occupational Safety and Health Administration (OSHA), National Institute for Occupational Safety and Health (NIOSH), and university environmental, health, and safety (EHS) rules procedures, labs will have to implement new protections to minimize the transmission of COVID-19 in workplaces. These protections can be expected to disrupt many of the processes in place before the COVID-19 pandemic, presenting a significant challenge to administrators who must manage a balance between preventing an outbreak and supporting productive research. This study aims to provide useful guidance for administrators, and the scientific background on which that guidance is based, helping to create a shared understanding among all involved in the research enterprise. This report does not contain a single comprehensive plan for ramping up research, but suggests widely applicable operational procedures, many of which are already being implemented in research environments.

Parts of this report uses the basic production number R as a conceptual object for thinking about mitigation measures. Conceptually, R can be thought of as a number measuring the dynamics of disease transmission for a particular community, and is defined as the average number of new infections spawned from an average infected person under the condition where nobody else in the community has immunity to the disease. It is important to point out that we do not know the baseline R , typically called R_0 , for the university or for the outside world. These values are functions of population demographics, the weather, the characteristics of the rooms and places where people congregate, and the nature of their interpersonal interactions. As an illustration of how difficult R_0 is to estimate for the community at large, retrospective epidemiological studies of the SARS-CoV-2 outbreak in China generally place the community-wide R_0 in the range of 2–6 [5]. When fully populated, college campuses are high-density, high-interaction environments, so one might expect the R_0 for a university to be larger than the community R_0 —though demographic elements such as age can have countervailing effects. Without prior knowledge of the university’s R_0 , administrators are left to make assumptions about how much they must do to prevent an outbreak, enforcing as many mitigation measures as they deem sensible. They can then increase or relax safety measures if they observe evidence of growing or tempered transmission. The effective R with mitigation measures in place is called R_{uni} for the university campus, and R_{comm} for the community in which

the campus is embedded. Using this concept, and taking the view that universities hold the responsibility for the safety and health of their population, we developed three principles to guide our thinking:

1. The operation of the university should not exacerbate the spread of the disease relative to the local community conditions. In terms of the widely used *basic reproduction number*, this can be expressed as $R_{\text{uni}} < R_{\text{comm}}$. Ideally, the university should aim to achieve a condition where $R_{\text{uni}} < 1$.
2. There is no one-size-fits-all solution. Each university, each unit, each lab group has different populations, activities, and infrastructures to which a general plan will need to be tailored.
3. Given how rapidly the scientific understanding of SARS-CoV-2 is evolving, it is necessary to undertake the research restart endeavor with a mindset of nimbly responding to changing circumstances. This requires collecting local data, keeping abreast of changing scientific assessments, analyzing the evolving situation, and adjusting operations in response to eventualities.

JASON has structured this report around eight questions a vice president for research (or equivalent) might ask.

1. What are the relevant characteristics of COVID-19?
2. How is the risk of airborne transmission?
3. What is the role of diagnostic testing?
4. How can health screening reduce spread?
5. How does one prevent a super-spreader event?
6. Is the campus an island from the community?
7. What operational policies are recommended?
8. How can institutions make risk-informed decisions going forward?

Each section of the report seeks to answer these questions through a narrative and concludes with findings and recommendations.

JASON concludes that a safe return to research can take place. The history of the 1918 flu pandemic (see Appendix A) offers a valuable lesson about mindset: quick and decisive action is needed at the outset, and perseverance so as not to relax restrictions too early. A ramp up will take months and requires careful adherence to rules and processes. Researchers will need to exercise patience and follow procedures that may hinder their productivity, but are ultimately necessary for public health. And all involved need to appreciate that a research restart may entail a rapid shutdown, as would be necessary if there is evidence that the incidence of infection is increasing.

2 What are the relevant characteristics of COVID-19?

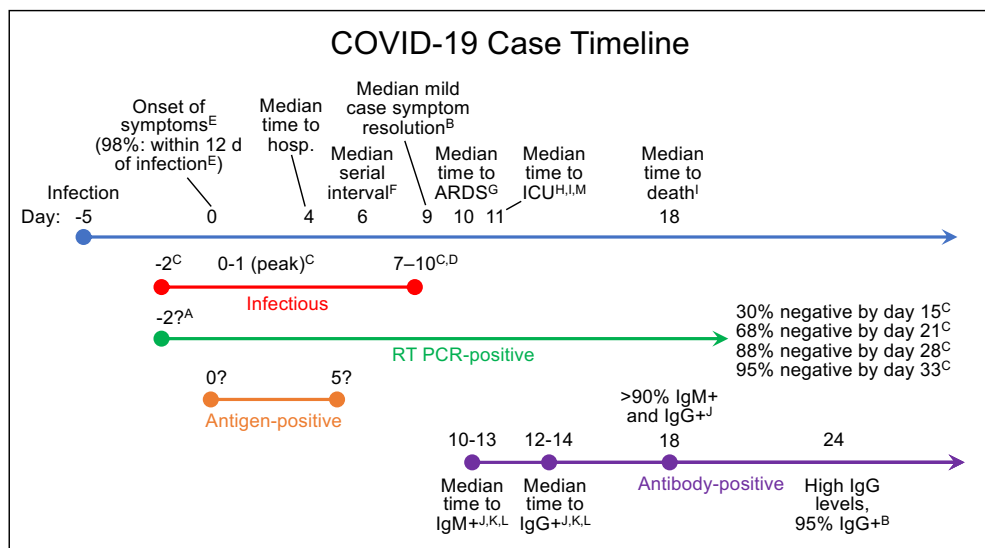


Figure 1: Timeline in days for the evolution of COVID-19 disease (blue), including the infectious period (red), PCR-positive period (green), antigen-positive (orange), and antibody-positive (purple). Symptoms onset occurs at day 0. Time points represent typical cases drawn from aggregated reports; actual event timing will reflect probability distributions centered approximately around these time points. ‘?’ indicates limited data or higher uncertainty. Mean serial interval is the average time until a secondary infection caused by the first becomes symptomatic. ARDS = acute respiratory distress syndrome; ICU = intensive care unit. Data are based on published reports as of 10 June 2020. References to figure labels are as follows: A[39], B[62], C[23], D[66], E[40], F[28], G[41], H[18], I[63], J[43], K[56], L[70], M[22].

Figure 1 presents a timeline of the current understanding of COVID-19 disease progression with an emphasis on the progress of cases that are sufficiently symptomatic to be diagnosed, and particularly those with respiratory symptoms. Several disease characteristics, particularly the onset of symptoms, the time period over which someone is infectious, and the time period over which various tests may register a positive result, are pertinent to designing a safe return to research. All the times in Figure 1 are for the median case; the population will exhibit some variation around these values and this variation needs to be considered in the design of policies. It

is important to note that individuals who are asymptomatic or only weakly symptomatic may also contribute to the spread of the disease.

Symptoms usually show about five days after infection, and 98% of cases show symptoms within 12 days of infection [40]. However, median infectiousness often precedes symptoms by about 2 days [23]. This presents a significant challenge for research restart. Screening for a fever or other gross symptoms leaves open the possibility that pre-symptomatic individuals are unknowingly spreading the disease on campus for several days. One study estimated that about half of all new infections originate from pre-symptomatic individuals [29].

Reverse transcription polymerase chain reaction (RT-PCR) testing is able to detect some infections during the pre-symptomatic period, with the probability of detection increasing towards symptoms onset [39]. It should also be noted that individuals can remain positive on a RT-PCR test for a significant period after they are no longer likely infectious because the body continues to shed non-infectious viral RNA for some time, with 68% becoming negative by day 28 and 95% being negative by day 33 [23]. For more on testing, see Section 4.

The time to symptoms onset is important for determining how long to quarantine individuals who have been exposed to a COVID-19 positive person and are thus potentially infectious. The standard quarantine period of 14 days was chosen by public-health officials before much information was available. More recent data suggests 98% of all infections will show symptoms within 12 days, permitting a shorter quarantine [40]. Even a quarantine of one week should catch about 80% of secondary infections, meeting the objective of $R_{\text{campus}} < 1$, assuming nearly all exposed persons can be identified. For more on estimating impacts on R , see Appendix B.2.

This timeline and associated brief summary reflect the state of science in early June 2020. The scientific community’s understanding of COVID-19 is rapidly evolving, and much data remains to be gathered and examined.

Findings

Finding: The disease characteristic most relevant to research restart is the high fraction of transmissions from pre-symptomatic individuals. Current data suggests the median time to the onset of infectiousness comes about

two days prior to the onset of symptoms.

Finding: A quarantine period of 7–14 days is appropriate for individuals who may have been exposed to SARS-CoV-2, with the lower end being sufficient to prevent a major outbreak on campus provided almost all exposed persons can be identified.

3 What is the risk of airborne transmission?

There is extensive evidence of airborne transmission of respiratory viral illness. Airborne transmission occurs when an infectious person breathes, speaks, eats, coughs, or sneezes, emitting small liquid particles that float in the air, and these particles subsequently come into contact with a susceptible person’s mucus membranes. It is extremely challenging to produce an estimate of the absolute risk of transmission for SARS-CoV-2 from airborne exposure, and we have not attempted one here. However, as a potentially illustrative example, Figure 2 shows an interpretation by Bueno de Mesquita et al. [7] of data from the largest human influenza challenge-transmission trial conducted to date, with 127 persons sharing poorly ventilated hotel rooms. The figure shows the estimated probability of influenza transmission as a function of cohabitation time. Note the exceptionally large variance between the 10th percentile estimate (dotted lines) and the 90th percentile estimate (dashed lines).

The suspended particles carrying virus exist in a continuum of sizes, from fractions of a wavelength of visible light to palpable droplets of spittle. Generally, the range is broken into two categories: *aerosols* that remain suspended for “long” periods of time, and *droplets* that fall to the ground quickly. The behavior depends on the precise environmental conditions and the choice of where to put the boundary is ultimately a subjective one.

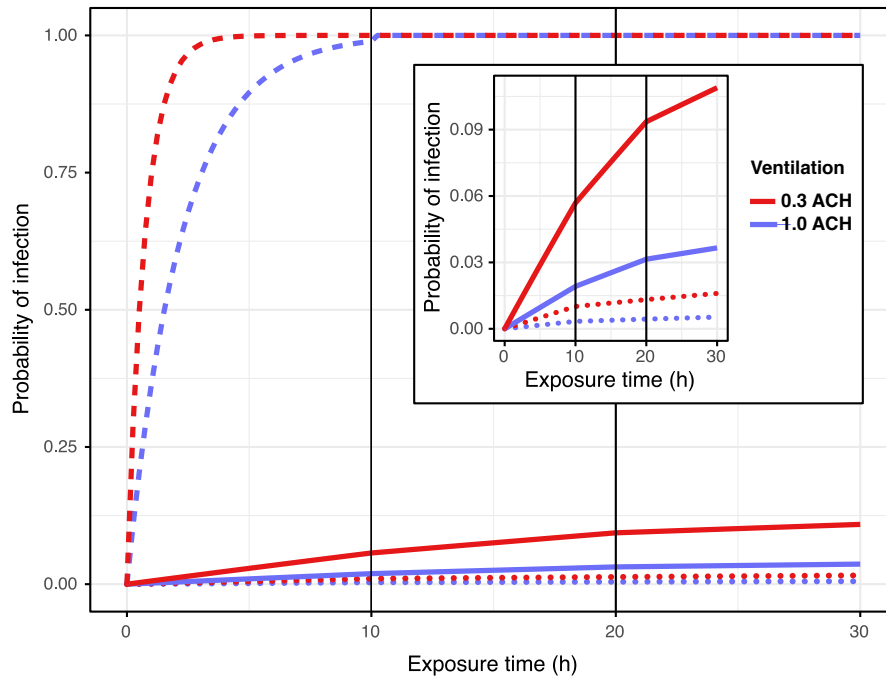


Figure 2: Probability of infection of a theoretical roommate in ultra-low-ventilation room (0.3 ACH, red) and low-ventilation room (1.0 ACH, blue). Solid line is the mean estimate, dashed line is the 90th percentile estimate, and dotted line is the 10th percentile estimate. Variation arises in large part because of the highly variable level of virus shed across different infectious persons (see Section 6). Inset graph is a zoom of the lower curves. ACH is the ventilation rate in air changes per hour. Figure adapted from [7].

3.1 Risk reduction

Strategies for reducing exposure to airborne particles include

1. physical distancing (e.g., the 6-foot rule);
2. wearing masks and eye protection;
3. reducing the number of people in a room;
4. modifications to ventilation systems.

In complement, strategies to reduce the release of airborne particles include:

5. wearing masks to reduce the emission of airborne droplets;

6. reducing particulate-generating activities such as speaking, singing, and eating in shared indoor spaces.

Each of these strategies is reviewed here.

3.1.1 Physical distancing

Physical distancing, such as the six-foot rule, emerges from the observation that a cough produces a jet of droplets and aerosols that travels about six feet into ambient air. Speech also produces a similar jet, only less powerful. The gas in these jets turns upwards because it is warm and humid and thus buoyant compared to ambient air. Small aerosols can be entrained in this upward plume [52]. Larger droplets, by contrast, separate from the plume under the influence of gravity and fall towards the ground. Most of this occurs within three to six feet [67], which is the origin of the six-foot rule. However, a small fraction of particles in the transition region will float at elevations where they can be easily inhaled, and may persist there for distances greater than six feet. These particles will tend to travel on the air currents in the room.

Quiet breathing entails conditions that relax somewhat the need for six feet of distance. Breathing produces essentially no large droplets and instead is limited to fine aerosols with less (albeit nonzero) viral load. For comparison, the volume of fine aerosol emitted by speaking is twenty times that of breathing, and the volume emitted over all particle size is hundreds of times that of breathing [49]. Thus, if two people must work closely for a short time, they should continue to wear a mask, not speak when in close proximity, and avoid being directly over another person to the extent possible.

Physical distancing can only partially mitigate the risk posed by the accumulation of aerosols and dispersion of aerosols in a room, which is also subject to room occupancy, airflow details, and activity intensity (discussed below). Thus, while physical distance helps avoid the great majority of the emitted viral load, especially outdoors, other measures are indicated to address the risk posed aerosols that build up in confined spaces.

3.1.2 Masks and eye protection

For large particles, masks and eye protection provide simple barriers to droplet projectiles from speech and cough jets. As particles become smaller, masks also filter inhaled air, becoming somewhat permeable to particles with diameters below about $10\ \mu\text{m}$. For more on the utility of masks, see Section 8.4.

3.1.3 Reducing the occupancy of the room

Under steady occupancy, the viral load in aerosols increases and levels off to a fixed value. In a room with well-mixed air, the exposure risk to any one individual is directly proportional to the number of infectious persons occupying the room. Existing research is insufficient to estimate the absolute risk of ongoing exposure to aerosol, although some authors have attempted conservative approximations [19, 3].

By contrast, the relative risk of adding more people to the room can be estimated. One way to characterize the relative risk is relative probability of at least one infection occurring from having n persons in the room relative to the probability of one infection from having 2 persons in the room. This relative risk factor is given by

$$P_{\text{rel}}(1, n) = \frac{1}{2}n(n-1). \quad (3-1)$$

The derivation for this factor and additional considerations for estimating the relative and marginal risks are discussed in Appendix G.

3.1.4 Modifications to HVAC systems

Heating, ventilation, and air-conditioning (HVAC) system modifications are one of the several ways to address the persistent risk of accumulated aerosols. A decision to modify the HVAC systems is not straightforward, however.

In the limit of a well-mixed room, increasing the provision of clean (outdoor or HEPA-filtered) air will help dilute the viral loads. However, increased airflow can also help move plumes of aerosols and droplets from speech, coughing, or eating across the room, increasing the risk of infection

to downstream individuals [42]. On balance, if speech is minimized, distancing measures are in force, and face masks are worn, then increasing the air supply in a room is probably beneficial—but a precise conclusion ultimately requires knowledge of the mask leakage rate, position of individuals, and air-flow patterns within the room. By contrast, in places where people eat, or if masks are not required, increased airflow may be counterproductive.

HVAC systems are usually designed either for displacement ventilation (where air inlets and outlets are at different elevations or on different sides of the room) or mixing ventilation (where air inlets are typically centered on the ceiling). Displacement ventilation provides superior air quality by displacing contaminated air in a bulk fashion. In the limit of perfect displacement ventilation, the rate of virus removal is roughly proportional to the air changes per hour (ACH), within the limits of normal ventilation rates where highly turbulent conditions are avoided.

In the United States, mixing ventilation is more common. With mixing ventilation, fine aerosols become diluted and distributed into the room. Under conditions of strong mixing (such as cool air entering from the ceiling) this can happen on a timescale of several minutes [9], but quiescent zones and streams remain possible. If aerosols from quiet breathing are homogeneously mixed into the room, the concentration in the room $C(t)$ can be described as by

$$\frac{dC(t)}{dt} = \frac{\mathcal{N}}{V_{\text{room}}} - \frac{C(t) \ln(2)}{\tau_{1/2}} - C(t) \cdot ACH, \quad (3-2)$$

where $\mathcal{N}$ is the constant emission rate of virus, V_{room} is the volume of the room, $\tau_{1/2}$ is the infectious half-life of SARS-CoV-2 at 1.1 hrs¹ [60]. In the limit as time $t \rightarrow \infty$, the equilibrium concentration becomes

$$C(\infty) = \frac{\mathcal{N}}{V \left(ACH + \frac{\ln(2)}{\tau_{1/2}} \right)}. \quad (3-3)$$

The effect of equation 3-3 is illustrated in Figure 3. More than half the benefit is achieved at 4 ACH, which is on the lower end of typical commercial spaces; and 90% of the benefit is achieved at 15 ACH, which is on the high end of what one might find in a laboratory space.

An alternative or complement to HVAC modifications would be to use

¹Measured at 65% relative humidity and 22 ± 1 °C.

true² HEPA-equipped air purifiers to remove aerosol loading. Such a filter has the potential to provide around 100 cubic feet per minute (CFM) for approximately 10 watts of power consumption. If the clean air from the unit displaces room air, it has the potential to provide localized areas of highly reduced aerosol loading while also helping to clean the rest of the air in the room. A 100 CFM unit in a $25 \times 25 \times 10$ foot room provides about 1 ACH of additional “fresh” air under the assumption of a well-mixed room.

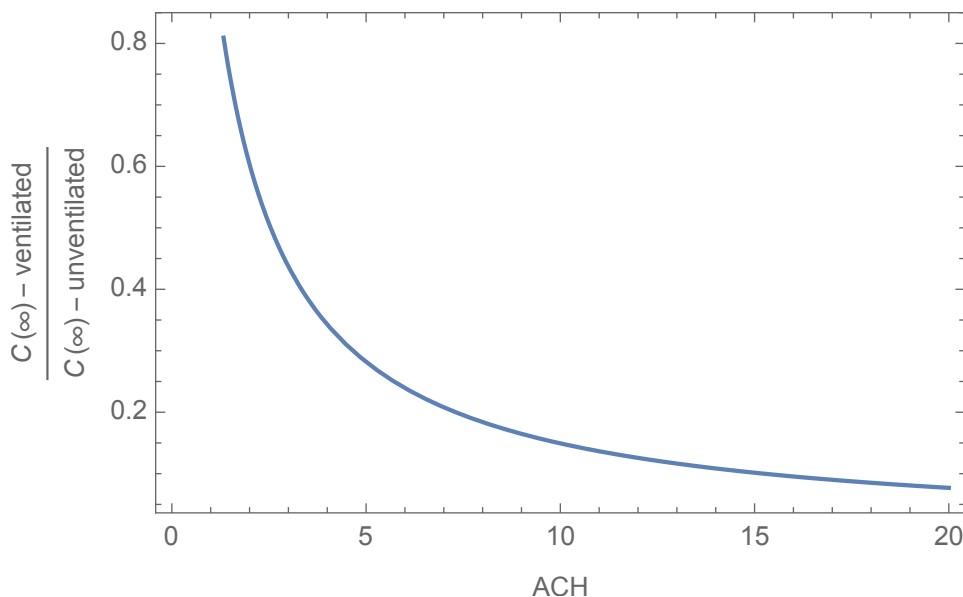


Figure 3: Diminishing benefit from increasing ventilation in the reduction of SARS-CoV-2 aerosol loads (exclusive of droplets).

Some buildings’ HVAC systems may circulate air between rooms in ways that are not apparent, allowing the spread of SARS-CoV-2 between rooms [19]. The campus facilities group must be consulted to ensure the air-flow is well understood and that the intended occupancy meets the capacity of the HVAC system.

²HEPA filters are high efficiency particulate air filters that meet consensus standards of the American Society for Testing and Materials (ASTM) or the American Society of Mechanical Engineers (ASME). These standards amount to filtering 99.97% of aerosol particles of $0.3 \mu\text{m}$ in diameter. Some inexpensive filters labeled ‘HEPA’ may not meet these standards.

3.1.5 Wearing masks as source control

The published literature finds that if an infected person wears a mask, it reduces to the risk to others [55, 48]. For example, Milton et al. [48] found surgical-style masks reduced influenza virus emissions by about 3 fold (95% CI 1.8–6.3), and reduced virus-bearing particles larger than 5 microns by 25 fold (95% CI 3.5–180). This suggests significant attenuation of larger droplets from speech is possible. However, the imperfect seal of masks may still allow significant leakage [57, 31, 13]. JASON was not able to ascertain whether face-mask wearing would fully counteract the additional risk posed by speech relative to quiet breathing.

3.1.6 Reducing droplet generating activities

Because the viral load is proportional to the volume of respiratory fluid, and because the volume of a particle scales as its diameter to the third power, the viral load shed by a person is mostly carried in the largest particles emitted. Importantly, essentially all particles produced when people are quietly breathing fall well into the category of aerosols [49, 34]. Whereas speaking, coughing, and sneezing not only produce more aerosols than breathing, they also produce larger droplets that contain substantially more virus than the aerosols. Speaking will increase the dose to others by at least 20 times, even when wearing a mask [49]. If larger droplets escape from masks, which seems possible because of imperfect seals, the amount of virus emitted by speaking can be many hundreds of times that of breathing [10, 68].

While speaking is an important part of academic work, minimizing unnecessary speech, and speech in close proximity to a person, can provide substantial benefits. Written text could provide a functional alternative for infrequent communication, especially for laboratory workers in tight spaces. Additionally, studies have found that that speaking “loudly” produced more than four times as many aerosols than speaking softly[2].³ Whispering, which does not involve movements of the vocal cords, produced about half the number of particles as “normal” speech, including a detectable reduction in larger

³The peak size also increased from 0.8 μm for soft speech to about 1 μm . It is unclear if this shift in size reflects changes in aerosol production, is an evaporation effect caused by reduced mixing time with ambient air, or is an effect of the geometry of the sampling instrument. If the effect relates to production, the increase in droplet diameter corresponds to another factor of two in emitted volume, suggesting that the total increase in emitted virus could be as much as eight times higher when speaking loudly.

particles around $2\text{ }\mu\text{m}$ [49]. The implications of this study are important for face mask usage. For example, if a face-mask attenuates 50% of outgoing particles (e.g., as might occur with a low-quality mask or respirator with a vent), then it would be better to speak softly and not wear a mask than to speak loudly with a mask. If a mask attenuates more than about 80% of outgoing droplets as implied for surgical masks by Milton et al. [48], then it is better to keep the mask on and speak at the lowest audible volume.

Findings and recommendations

Finding: Physical distancing does not fully address the risk posed by aerosols in a room, which become dispersed in the room’s air.

Finding: Masks reduce the risk for the wearer as well as the risk posed by the wearer to others. However, masks that force people to speak loudly may be counterproductive unless their outgoing filtration efficiency is high.

Finding: The relative risk of adding more persons to a room goes roughly as the square of the number of persons. Thus, doubling occupancy increases the risk by a factor of four.

Finding: The impact of increasing airflow on disease transmission depends on the activities occurring in and the layout of the room. When combined with masks, distancing, and minimization of loud speaking, increasing airflow is probably beneficial.

Finding: For mixing ventilation, there are diminishing returns to higher airflow. More than 50% of possible gains are had at 4 air changes per hour (ACH), 90% of gains at 15 ACH.

Finding: Speech sheds at least 20 times more virus than breathing on a minute-by-minute basis, even if wearing a mask. Viral shedding may be up to several-hundred fold depending on conditions.

Recommendation: Reduce unnecessary speech communication by using text through text messaging, scratch pads, white boards, or a dedicated computer and flat-panel display. If you must speak, speak softly.

Finding: Require all occupants of indoor rooms to wear masks. Respirators

with easy-exhale valves provide limited source control, but can be covered by an additional surgical mask.

Recommendation: Ensure room ventilation meets minimum standards. If your institution chooses not to require masks for all occupants, consider a study of how airflow patterns may transport plumes and larger droplets before increasing airflow.

Recommendation: Distribute personnel equally across rooms to minimize the number of persons per room; and equally within rooms to reduce the transmission of larger droplets.

4 What is the role of diagnostic testing?

There has been extensive discussion of “testing” as a potential requirement for restarting universities. Much of this discussion is ambiguous about the kind of testing or its value. There are many different kinds of diagnostic tests for SARS-CoV-2 that can be employed to serve substantially different ends, each with a different cost, logistics burden, error rate (false positives and negatives), turn-around time, and level of patient discomfort. This section discusses how these tests might be used as part of a research restart.

4.1 The kinds of tests

Testing for the presence (or past presence) of SARS-CoV-2 can be broken into three categories based on the target of detection:

1. Viral RNA testing;
2. Viral antigen testing;
3. Antibody testing.

4.1.1 Viral RNA testing

Viral RNA testing seeks to detect the genetic material of SARS-CoV-2 and can be used to identify residues of genetic code from the virus in tissues, fluids, surfaces, air, etc. Essentially all of these tests operate by using enzymes that “amplify” or replicate the genetic material present in the sample. Quantitative versions can detect not only the presence of viral RNA but also estimate the amount of RNA present in the original sample—offering a proxy for viral titer and infectiousness. Importantly, however, none of these tests determine whether the RNA detected came from viable virus particles capable of infecting cells, or merely the debris of dead virus. Thus, these tests can indicate the presence of RNA long after the virus has ceased to be infectious. There are several different methods of RNA amplification, each offering different features as they relate to a diagnostic setting.

PCR—PCR, or more precisely, reverse-transcriptase polymerase chain reaction (RT-PCR) is an enzyme-based reaction that amplifies genetic code

segments in a series of thermal cycles. The test is extremely sensitive, but it requires well-controlled conditions and either trained personnel or robots to carry out. PCR is currently the diagnostic standard in the United States, with many small variations in the kits supplied by different commercial vendors.

Initially, PCR tests were only validated for highly invasive nasopharyngeal swabs that require a trained person to collect and were very uncomfortable for patients. Because nasopharyngeal swabbing tends to induce sneezing, medical staff also need to wear full personal protective equipment. At the time of this writing, some labs have begun to offer PCR processing of less invasive swabs taken from the anterior nares (nostrils), and there are now efforts to standardize a saliva-based test that can use samples collected by the patient.

Despite PCR’s exquisite sensitivity, the overall efficiency of the sample-collection procedure and the variable expression of virus in human tissues at different stages of the disease can result in false negatives. A recent analysis by Kucirka et al. [39] found that over the 4 days between infection and the typical time of symptoms onset (day 5), the probability of a false-negative in nasopharyngeal samples went from 100% (95% CI, 100% to 100%) on day 1, to 67% (CI, 27% to 94%) on day 4. In other words, on days prior to a person showing symptoms, the test is more likely to give a false result than a correct one. On the day of symptoms onset, the median false-negative rate was down to 38% (CI, 18% to 65%). Despite this low probability of detecting an infection before symptoms, the test may be better at detecting infectiousness to others. This is because infectiousness is a function of the person’s expressed viral titer, with higher titers producing lower false negatives. However, to our knowledge there are no studies attempting to quantify the predictive value of a PCR test on infectiousness. In contrast to false negatives, the test tends to have a very low false-positive rate [61, 33].

With the exception of a few hospital-grade robotic systems, most PCR-based tests require sample storage and transport to a certified laboratory capable of performing the needed RNA extraction, concentration, and PCR-based amplification. As such, these tests typically have long (12 hour to 2 day) turnaround times and are fairly expensive (\$50–\$100/test when processed at scale). This means PCR tests are suitable as a diagnostic of persons with suspected infection or exposure to the virus, but PCR tests cannot currently be performed quickly and cheaply “at the front door.”

LAMP—Reverse-transcriptase loop-mediated isothermal amplification (RT-

LAMP) is another amplification system that differentiates itself by using a simplified single heating cycle to process the sample. This circumvents much of the complication associated with PCR-based testing. In combination with chemical indicators or dyes, this test can be designed to yield a yes/no answer by simply observing a color change in a small plastic test tube. The time to process a sample is a few minutes, followed by about 30–60 minutes of waiting. The complication and cost of running the test is thus substantially reduced. With the amplification steps simplified, sample preparation steps become the primary burden. Some test designers are discarding sample preparation steps to improve ease of use at the cost of reducing the sensitivity to 1/10 its nominal value. If the test can be validated using saliva or self-sampled anterior nares swabs, then virtually all the steps requiring trained staff can be eliminated. With all steps simplified, commercially provided test kits could in principle be easy enough to perform virtually anywhere with minimal training and no more equipment than an electric heat block. However, FDA approved versions are unlikely to be available except to certified diagnostic laboratories unless the FDA deems the test robust enough to make a special exception for SARS-CoV-2.

The trade-off for the simplicity of LAMP tests is a reduction in sensitivity, resulting in higher false-negative rates. In addition, the genetic segments identified by LAMP test are shorter than in PCR tests, which means that they tend to be less specific, resulting in higher false positives than PCR. A non-clinical laboratory environment is also more likely to introduce variables such as poor temperature control, further increasing both false positive and false negative results. Based on informal pre-clinical data from test manufacturers, it appears that the error rates for LAMP tests are still likely to be low enough to serve as a useful detector of infectious persons.

4.1.2 Viral antigen testing

Viral antigen tests directly detect the proteins of the virus rather than its RNA. Samples are diluted and specially made antibodies bind with the viral proteins, usually carrying with them some dye or fluorescent indicator up a lateral-flow strip. These tests tend to be highly specific, but are not as specific as RNA-based tests, meaning their false positive rates are higher but still acceptable.

The FDA has approved one laboratory-grade antigen test at the time of writing. That test uses a dedicated machine to read the test outcome

and is designed to be used only in laboratory settings. The approved test has about 80% of the sensitivity of PCR tests and the specificity is such that it will give a positive result for SARS-CoV-1 and SARS-CoV-2 [53]. By convention, antigen tests are usually read in 15 minutes to maximize their sensitivity, but the most highly infectious persons will show positive results in about one minute. They are also relatively inexpensive. The cost of an antigen-test cassette can be less than \$10. Some manufacturers are taking advantage of this fact by working to win FDA approval for viral antigen tests that do not require a machine and which have as their purpose a near-instant detection of the most infectious persons. While these tests will have lower sensitivities, the extreme speed and simplicity means they may have a role to play in identifying the most infectious persons and super-spreaders “at the front door.”

4.1.3 Antibody testing

Also called serological tests, antibody tests aim to determine if a person has been previously exposed to SARS-CoV-2 virus by detection of antibodies produced by the human immune system. Unlike viral RNA and antigen tests, antibody tests do not detect the presence of virus. It is likely that these antibodies will confer some protection against reinfection, and it is possible that the protection conferred from a strong immune response may last for a period of over a year [4, 65]. However, it is not yet known what level of antibody titer is needed to prevent re-infection, or how long a person who has had a mild infection might be immune to reinfection. Many antibody tests are not sensitive to antibody titer, giving only a qualitative measure of the presence of antibodies. Thus, it is not yet clear how antibody test results should be interpreted, or how those results are actionable from a public-health perspective. Under particular situations, knowledge of plausible antibody presence might help optimize how people are deployed in the battle against COVID-19 (e.g., tasking plausibly immune doctors and nurses to treat COVID-infected patients over those who are immune naive). But given the large uncertainties, it is not evident that antibody tests should be pursued by groups not needing to manage particular exposure risk.

A large number of antibody tests were approved by the FDA under Emergency Use Authorization; some have now lost that authorization. Many of the approved tests appear to have high false-positive and false-negative rates—making an already difficult interpretation problem even more challenging. In general, it seems wise to avoid creating a situation where anti-

body tests create a privileged class of persons who are deemed immune and thus not subject to the same restrictions or protections as others. Not only might unequal application of protection rules complicate compliance, it may have the perverse effect of moral hazard, creating an incentive to purposefully increase one's risk of exposure in the hope of moving to the privileged class.

4.2 The functions of testing

We discuss here three potential functions of testing, the qualities tests should have, and the role these functions might practically play in research restart. The three functions are: diagnosis, screening, and monitoring.

4.2.1 Diagnosis

Tests are essential to the clinical diagnosis of potentially exposed or symptomatic persons. A positive test justifies follow-up patient referrals as well as recommendations for isolation as a precaution to protect the community. It is essential for administrators to ensure access to medical services with adequate diagnostic-testing capability. Ideally, these tests must have high sensitivity and low error rates. Of the technologies described in Section 4.1, PCR tests are the most suitable. Factors like cost and convenience may be deemed less important here. The same is true of test turn-around time, if additional transmission precautions are taken in response to symptoms prior to one (or two) negative tests. Low false-negative rates are particularly important to ensuring infectious persons are not incorrectly given a clear pass. It is possible to substantially reduce the impact of modest false-negative rates by testing persons with a small (e.g., 1 day) wait in between—provided the test does not also have a high false-positive rate.

4.2.2 Screening

Screening implies testing the entire population or, a defined subset, on a regular basis. This places greater emphasis on cost, comfort, and convenience. A saliva-based test would be ideal because it is non-invasive. For screening, it is also crucial the test results be had and become actionable within

a short period of time. If a test returns a result after one day, for example, an infectious person will have been spreading the disease during that time. If the results take longer than two days, an infectious person is likely to develop symptoms on roughly the same timescale as the test result, rendering most of the tests moot (see Section 2 for a discussion of these timelines). Finally, as the goal is to reduce the spread of the disease on campus, it is more important to identify those who are highly infectious than those who are infected but only mildly infectious. An especially valuable aspect of a test-based screening would be its ability to detect super spreaders (see Section 6). All things considered, this suggests there are incentives for trading away some test sensitivity for speed, convenience, and lower cost.

Both RNA and antigen tests can be used for screening. PCR tends to be slow and expensive, LAMP considerably less so. Antigen tests have the potential to be particularly fast and cheap because they usually do not need wet chemistry to process a sample, and the most infectious individuals will show a positive result in about one minute. Thus, antigen tests have the potential to enable testing “at the front door.” In Appendix D, we show that such a real-time test performed daily before the start of the working day reduces the expected exposure time for coworkers by approximately a factor of 5 relative to a laboratory-based diagnostic in which samples are taken before going home and the results are processed in a lab overnight with the result returned before work the next morning. Although the sensitivity of rapid screening technologies is lower, the reduced exposure time for coworkers enabled by the rapid read-out provides substantial compensation, especially when considering that rapid tests are likely to be sensitive enough to catch the most infectious persons, such as super spreaders. On June 16, 2020, the FDA announced that it would now consider approving test implementations for screening purposes absent a prescription from a physician. Test makers can now submit requests for emergency-use authorization of tests specifically designed for screening. However, at the time of writing, no laboratory-free rapid screen has been approved by the FDA or by state-government officials.

In Appendix B we derive equations that model the impact a generalized, continuous screening program might have on controlling the transmission of disease using highly simplified compartment models. Section B.1.2 give additional considerations for when the screen is a virus test. These findings result in equation 4-4, below, which can be used to help a university administrator estimate the benefit of a testing program relative to other risk-reduction measures. In this equation, the campus’ initial transmission environment is characterized by the basic reproduction number R_0 , which is equal to the average number of secondary infections caused by an average infectious person

if the entire campus population had no immunity.⁴ The effect of testing can then be expressed as a multiplier on that parameter, yielding a new basic reproduction number

$$R = R_0 \times \frac{1}{1 + r_t(1 - f_t)/\gamma}, \quad (4-4)$$

where $0 < f_t < 1$ is the test's false-negative rate, r_t is the continuous testing rate in tests per unit time, and γ is the rate at which ill persons are removed from the susceptible population absent testing, and is equal to the inverse of their average infectious period. Realistically, γ will fall between two extremes: if ill persons circulate in the population throughout their illness with no regard for others, $\gamma \approx 1/(6 \text{ days})$; but if perfectly effective health screening removes people as soon as they show symptoms, the serial interval (see Section 2) suggests $\gamma \approx 1/(1 \text{ days})$. We judge that health screening is valuable but will be far from perfectly effective (see Section 5).

Equation 4-4 makes clear the benefit of testing frequently (causing r_t to be large). Such a proposition is expensive and logistically burdensome. One approach to minimizing the burden would be to consider targeted testing of only those persons judged more likely to be infected. Such persons those who were recently in contact with an infected person as identified by contact tracing, or individuals engaging in riskier activities like commuting by public transport, air travel, or living in a fraternity. In Appendix B.1, we give an example where combining less frequent but targeted testing with regular symptoms attestation leads to the greatest benefit, reducing R_0 by over 50%.

Another approach to minimizing the burden of testing is to consider pooled testing, in which multiple samples are combined and the test is processed to see if there is evidence of virus among the whole group. If a group of samples comes back positive, each member of the group must be individually re-tested (at least once, given false-negative rates) to eventually identify the infected person—and until that happens, all members of the group should quarantine. This strategy makes the most sense for expensive laboratory tests, like PCR. For realistic scenarios, pooled testing may reduce the costs of the laboratory step by about a factor of five; but sample collection costs are unchanged, and sample and handling cost may go up slightly because of the need to combine samples. The consequence of pooling samples, however,

⁴The basic production number is a conceptual object that includes effects from the disease itself, the characteristics of the susceptible and infected populations, the weather, the characteristics of places where people congregate, and the nature of their interactions. There is no true value of R_0 for a disease class of persons. Estimates for the population-weighted reproduction number in Wuhan, China have ranged from 2 to 9.

is a reduction in test sensitivity, causing a higher false-negative rate (f_t in equation 4-4). The change in sensitivity is not proportional to the number of tests being processed, however. This number would need to be determined for the particular test under consideration before one could evaluate the impact of reduced sensitivity on the false-negative rate. The benefit of pooled testing is also severely curtailed by false positives, which can rapidly lead to excess re-testing, washing out all cost savings. The study of pooled testing is a mature subject; Appendix E gives a more in-depth but still incomplete discussion by way of a simple example.

4.2.3 Monitoring

In principle, testing can identify current levels of active infection or the prevalence of previous exposure—but this requires designing a testing program that has sufficient statistical sampling. When the prevalence is low, false positive rates become particularly problematic as small changes in the prevalence may be washed out by false-positive noise. For example, consider a situation where the fraction of infected-but-asymptomatic people in the community is 2 per thousand ($p = 0.002$). A medium-sized university restarting at 25% of normal density may have 2,500 researchers on campus, ~ 5 of whom will be infected and asymptomatic. If the university has the capacity to test n people per day, the probability of selecting k infected persons in each day's test is

$$P(k) = \binom{n}{k} 0.002^k (1 - 0.002)^{n-k}. \quad (4-5)$$

If the university can test $n = 500$ per day (thus testing everybody once per week), the probability of getting no infected persons ($k = 0$) is 37%, and the probability of getting 1 infected person ($k = 1$) is also 37%. If the false positive rate is 1% then, on average, each day of testing will produce 5 false-positive tests, swamping the true positive detections. Retesting can help mitigate this, but if the false-negative rate is also high, then retesting has limited value. For example, if the false-negative rate of presymptomatic infections is 80%, in line with estimates for nasopharyngeal PCR reported by Kucirka et al. [39], then the probability of detecting each true positive person in the day's draw drops to 4% per person. Thus, virtually all (96%) of truly infected persons will not be detected, and virtually all positive results will be erroneous. A further illustration of the impact of false positives is given in Appendix H. Guidelines for determining the minimum number of people that must be tested to achieve a finding with a given confidence is given in Humphry et al. [32].

A more cost effective alternative to monitoring the incidence on campus would be to monitor the rate at which people are diagnosed based on self-reported symptoms, and to ensure clinical diagnosis occurs through a health screening program that directs symptomatic individuals to get tested. Such a program is discussed in Section 5.

Findings and recommendations

Finding: Antibody testing has, at best, a limited role to play in a research restart program.

Finding: The more rapidly a test result can be had and the results acted upon to quarantine of an infected person, the more useful testing will be.

Finding: Antigen testing may be attractive for “testing at the front door,” and detecting super spreaders, but such tests are not yet FDA approved.

Finding: Testing against asymptomatic populations requires a test with a false-positive rate well below the disease prevalence to produce useful insights.

Finding: Re-testing to reduce false positives also increases false negatives, and vice-versa. The interaction can have a devastating effect on the overall efficacy of using testing for screening asymptomatic populations.

Finding: False positives and false negatives make monitoring the state of the campus through testing difficult. Monitoring may be best done by tracking the rate of clinically diagnosed cases arising from symptomatic populations, and comparing the result to the larger community rate. A health screening program can help ensure good coverage.

Finding: Testing only the populations with the highest risk of infection substantially improves performance for any given investment in testing. Contact tracing is one method of identifying at risk individuals. Other factors include living situation, commuting habits, and immune status.

Finding: While pooled testing can be used to reduce the cost of testing, one must consider the effect of false positives and false negatives. Ultimately the gains from pooling may be small compared to a less-expensive and less-sensitive test.

Recommendation: Before deciding to test, perform a statistical analysis to understand how the false-positive and false-negative rates of the tests available to you impact your ability to meet your objectives.

Recommendation: Monitor developments of testing technologies, especially those that may produce results rapidly without needing a diagnostic laboratory.

5 How can health screening reduce spread?

Screening is the general strategy of identifying infectious persons to remove them from circulation and attenuate the spread of the disease. The value of different screens depends on how early in the course of the disease they can be effectuated, and the false-negative rate of the screen.

Testing for virus, described in Section 4 is the most direct method for identifying infectious individuals, but many currently approved tests are expensive, logistically burdensome, and subject to high error rates. This section will focus on health screening, either through a survey tool (attestation) that asks people to report on symptoms they may experience, or by direct physiological sensing.

Health screens detect effects of the immune response to disease. They are, therefore, second-order measures of a person's infectiousness. A critical concern with relying solely on health screening is that pre-symptomatic individuals are believed to be the source of a large fraction of secondary infections, with one study placing the number at roughly half of all new infections [12]. Nonetheless, health screening can be made relatively simple and low cost. A daily attestation of symptoms via smart-phone app would be one example. Taking every person's temperature at the front door is another. It is the frequency and ease of implementation that allow these relatively insensitive screens to contribute significantly to limiting the spread of disease. Voluntary reporting of symptoms or the reporting of no symptoms constitutes the collection of health data that must be collected and stored securely. Campus general counsel and health services will have to give guidance to ensure health information is correctly handled and viewed only by the appropriate people.

The symptoms associated with COVID-19, with incidence frequency and standard errors, include:

1. Fever (0.64 ± 0.030) [41, 44, 66]
2. Sinus pain (0.50 ± 0.18) [66]
3. Cough (0.46 ± 0.032) [41, 44, 66]
4. Reduced or altered sense of smell or taste (0.44 ± 0.17) [66]
5. Expectoration (0.32 ± 0.036) [44]

6. Stuffy nose (0.25 ± 0.15) [66]
7. Chills (0.18 ± 0.044) [41]
8. Fatigue (0.18 ± 0.025) [41, 44]
9. Sore throat (0.13 ± 0.039) [41]
10. Headache (0.13 ± 0.037) [41, 66]
11. Difficulty breathing (0.11 ± 0.034) [41, 66]
12. Joint or muscle pain (0.099 ± 0.023) [44, 66]
13. Diarrhea (0.056 ± 0.015) [41, 44, 66]
14. Vomiting (0.026 ± 0.018) [41]

To illustrate the limited sensitivity of these screens, consider that only 64% of positive cases report fever *at any time* during the course of the disease, and fever may also not be the first symptom of the disease [44]. Thus screening individuals for fever alone should be expected to miss at minimum 1/3 of symptomatic cases. Algorithms that combine multiple reported symptoms have been shown capable of retrospectively identifying 65% of positive cases [47]. Published algorithms are a starting point for using health-screening data, but more sophisticated approaches tailored to the campus environment are possible.

For the purpose of screening on campus, the fact that symptoms are ambiguous as to origin creates false positives; these false positives will increase during the flu, allergy, and cold seasons. Tracking population-averaged trends in the campus community allows for the sensitivity attributed to each symptom to be adjusted on a continuous basis, reducing both false positives and false negatives. Stress or chronic conditions among certain individuals will tend to render algorithms tuned for the general population less predictive for those individuals. Algorithms can be made adaptive to each individual's baseline state, compensating for person-to-person variation and increasing sensitivity. By contrast, the fact that symptoms may be mild or not present simultaneously, and that individuals may engage in deceptions because of a desire to work, will increase the false-negative rate. In sum, it is not yet clear to what extent screening for multiple symptoms will increase a university's ability to identify infected individuals and a high false-negative rate should be assumed at the outset so as to not overestimate the impact of health

screens. An epidemiological compartment model for estimating the impact is given in Appendix B.1.

Even assuming a high false-negative rate, the fact that symptom attestation is low cost and can be used frequently makes it a very useful component of a university’s toolkit. As shown in Appendix B.1, when an 80% false-negative rate is assumed for symptom attestation, and the attestation occurs only every other day, R_0 can still be reduced by 38%. To compare, if the rate of testing is once every 14 days for each person, and the tests have an optimistic 25% false-negative rate, R_0 is reduced by only 25%. The fact that symptom attestation can occur frequently substantially compensates for its low sensitivity.

In addition to yes/no reporting of symptoms, university researchers may wish to offer an opt-in version of the survey that allows willing participants to provide more detailed information. Useful data may include daily reports of one’s body temperature (taken with the identical thermometer), overnight respiratory rate (from a wearable device), or blood-oxygen saturation. Research on habits could also be useful, such as reporting instances of high-contact events such as grocery shopping or public transport. It is important, however, that most of these more burdensome questions of uncertain value be voluntary, so as not to reduce compliance with the primary survey.

Finally, in addition to symptom reporting, the medical team should consider adding, on occasion, other questions to the daily attestation. For example, a semiweekly numerical evaluation by each researcher about how safe they feel in their lab; and biweekly questions about whether individuals feel increasing stress. Such data could reveal problems with conditions on campus, or unsafe laboratories, both of which could undermine efforts to prevent COVID-19 infections.

Findings and recommendations

Finding: Modification of the standard SIR model to account for screening illustrates that a health attestation can substantially slow disease transmission. The high frequency of the screen and ability to act quickly compensates for the expected high number of false negatives.

Finding: The more rapidly screening results in the quarantine of infected persons, the more effective screening is.

Finding: There will be significant variation in the baseline symptom rate across the population because of chronic conditions. Additionally, changes in COVID-19 infection rates and non-COVID illness rates both change the predictive value of symptoms in opposing ways.

Recommendation: Develop a procedure for daily health screens, such as an attestation of symptoms before arriving at work. A smartphone app for attestation offers the valuable opportunity to automatically determine if this threshold is reached and thus act instantly on the information.

Recommendation: The university should mandate that a member of the university community displaying symptoms typically associated with common disease (like a cold) should not report to work. If in addition the individual complains of symptoms associated with COVID-19 or influenza-like illness, a diagnostic test for SARS-CoV-2 should be performed.

Recommendation: Develop a capacity to continually adjust the thresholds at which stay-home and testing-referral decisions are made. Adjustments will be needed based on the prevalence of COVID-19 in the community, changes in the understanding of the disease, and seasonal illness like influenza. Ideally, algorithms should automatically adjust to account for person-to-person variation.

Recommendation: Assume a high false-negative rate for symptom attestation when planning a restart.

6 How does one prevent a super-spreader event?

The concept of super-spreading is well known in epidemiology: it is the propensity of a single infected individual to infect a larger-than-average number of people. The effect arises from a combination of biological, behavioral, and environmental variables, all of which influence transmission [8].

This phenomenon is often associated with the 20/80 rule: 20% of the host population contributes at least 80% of the net transmission potential (as measured by the basic reproduction number, R .) The rule implies that control programs targeted at the core 20% group are potentially highly effective. Conversely, programs that fail to reach most of this group will be less effective in reducing levels of infection in the population as a whole [64]. In the case of SARS-CoV-2, this ratio may be closer to 5/95, as shown below in Figure 5 and the associated discussion below.

Are there distinguishing aspects of infection and transmission that might identify super-spreaders or circumstances leading to super-spreading events? One element of such identification, as discussed below, is that within the population of infected persons, some may carry a viral load 10^3 to 10^5 times higher than the modal case. Importantly, given the data available, this appears linked to the stage of the infection. In addition, some infected persons may have a higher potential to spread SARS-CoV-2 through the ways in which they speak, cough, or exhibit other personal characteristics. Comorbidities may increase their ability to transmit the virus, as noted below. Their interaction with their environment can contribute to increased transmission if the infected engage in high risk behavior such as attending large gatherings, riding on public transport, not wearing a mask or donning it improperly. The environment may contribute through poorly-designed ventilation.

There are large differences in the number of aerosol particles produced during breathing among different people (coefficient of variation around 1.2; see Section 3 for a discussion of aerosols and droplets). This variation amounts to roughly a factor of 100 between the 95th percentile emitter and the 5th percentile emitter and is believed to arise because of variations in mucus surface tension in the lung [15]. Variation in droplets produced during speech, which carry much more virus but are also substantially attenuated by both interpersonal distance and mask wearing, was recently found to vary by a factor of about 4 between the 95th percentile emitter and the 5th percentile emitter [2].

While variations in aerosol and droplet production are significant, they are small compared to the much larger variation in viral titers observed across COVID-19-infected individuals. In particular, differences of up to $8 \log_{10}$ (10^8) in viral load between individuals have been observed—a million times larger than the variation in aerosol production [1]. Importantly, this variation is not variation in the peak titer, but variation in titers estimated from virus tests at whatever time they were performed. The variation in peak titer will be smaller.

Person-to-person variation in immune-system strength appears to contribute partially to the variation in observed titers. One study found that peak viral titer in patients with certain comorbidities was, on average, greater by a factor of 100 [59]. Viral titer also appears to increase with patient age, which is correlated with slower immune response [59, 36, 28]. In particular, To et al. [59] estimates titers increase by a multiple of 7.5 per each decade of age. Independently, aerosol emissions appear to double between age 20 and 40 [35, 2]. Combining the two effects, it appears plausible that the typical 50 year-old professor could shed 1000 times more virus in the form of aerosols than the typical 20 year-old student at the peak of infectivity.

The data from Jones et al. [36] can help illustrate why individuals with high viral titers may dominate the spread of the COVID-19. The study reports the results of 3,712 patients who tested positive for SARS-CoV-2. Figure 4 reproduces the data showing how many samples fell into each of a set of bins of estimated viral load on the swab. The bins range from 10^4 (limit of detection) to 10^{12} copies/swab. Some of that variation can be attributed to infection age (changes in titer over the course of the disease), and some variation to variables associated with the collection of patient samples.

Multiplying the abscissa and ordinate at each point in Figure 4 yields the total virus in the bin, and normalizing over the sum of virus in all bins, yields the relative contribution of each bin to the total sampled virus, as shown in Figure 5. Assuming the data represents a snapshot of the infected population, Figure 5 shows that at any point in time the majority of the virus being shed is being shed by a small group of people. Those in the last seven bins comprise only about 5% of the population but contribute about 90% of the total virus being shed.⁵

⁵To be clear, this particular analysis is uncorrelated with other identifying factors and thus does not by itself say whether those in the most infectious bins are identifiable in some way, or whether the “average” person will pass through one of these bins briefly during the course of the disease.

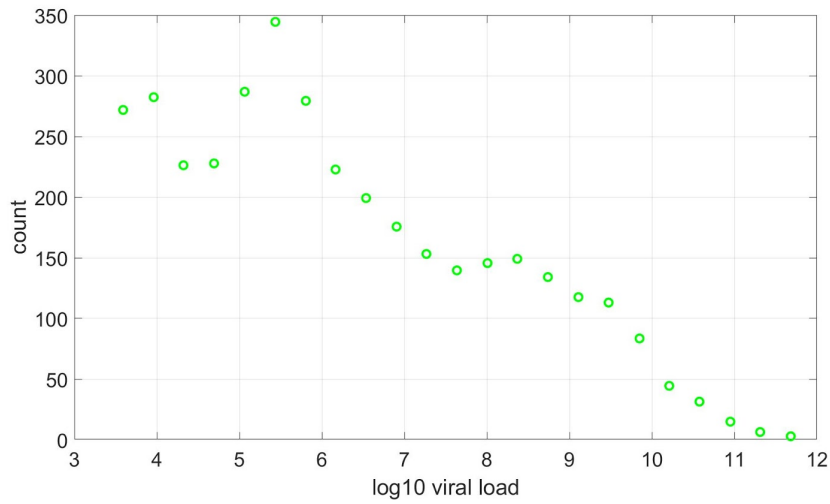


Figure 4: “Histogram of viral loads: The plot shows the frequency distribution of 3,712 values of patient SARS-CoV-2 (logarithm base 10) viral load, estimated from real-time RT-PCR Ct values... The sharp drop on the left side of the distribution is due to RT-PCR sensitivity and the limit on the cycles.” Caption and data extracted from a figure in [36]

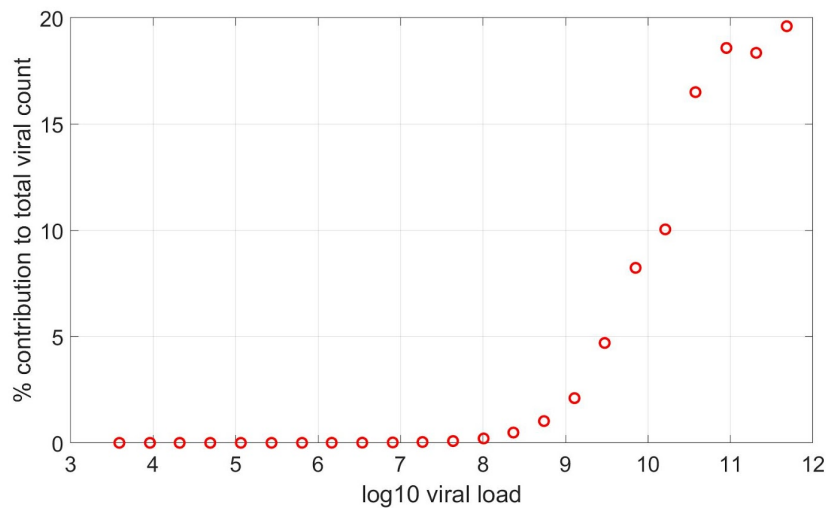


Figure 5: Total amount of virus present in each bin of Figure 4. This plot helps visualize how much virus is contributed by patients in each bin to the total amount of virus present in the sampled population.

As a final consideration, it is notable that viral titers appear to be highest around the time, and possibly just before, symptoms onset. While not yet robustly demonstrated, the inference that viral loads are at least as high just before symptom onset as they are when first measured is apparent from Figure 6. The analysis of Kim et al. [38] supports this assessment: “In sensitivity analysis, using the same estimating procedure but holding constant the start of infectiousness from 1 to 7 days before symptom onset, infectiousness was shown to peak at 2 days before symptom onset.” This suggests that super-spreaders are likely to be present among the pre-symptomatic population, making the identification of super-spreaders a special challenge.

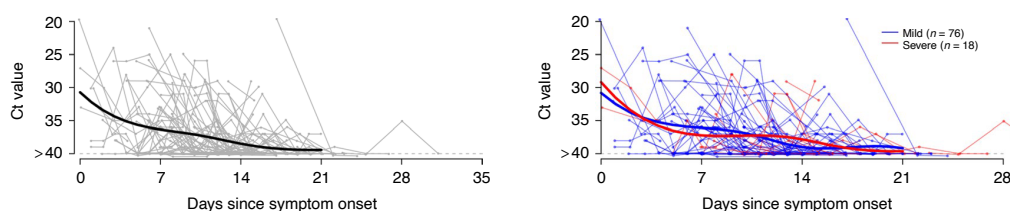


Figure 6: Figure and caption reproduced from [28]. “Viral load (threshold cycle (Ct) values) detected by RT-PCR in throat swabs from patients infected with SARS-CoV-2 (N=94), overall and stratified by disease severity... The detection limit was Ct=40, which was used to indicate negative samples. The thick lines show the trend in viral load, using smoothing splines. We added some noise to the data points to avoid overlaps.”

Findings and recommendations

Finding: Super spreaders will contribute disproportionately to the number of secondary infections; they cannot be neglected in the design of protective measures.

Finding: In general, data suggest older persons are both more infectious to others and more at risk of infection.

Finding: Super-spreaders are likely to be pre-symptomatic individuals with high viral titers.

Finding: If super-spreaders are pre-symptomatic individuals, health screens that depend on symptoms will prove less effective than expected based on model results obtained using the average infectious case.

Finding: Virus testing (both antigen and RNA tests) has the potential to identify super-spreaders with exceptionally high viral titers.

Finding: A symptomatic person with an exceptionally high viral titer at diagnosis has a good chance of having been a super-spreader. Immediate quarantine followed by contact tracing for these individuals will be especially helpful in attenuating the rapid spread of the disease.

Recommendation: If the technology becomes accessible, consider a program to screen by viral testing all asymptomatic people on a regular basis, ideally before work and perhaps twice per day. Such tests do not need to be sensitive to low viral titers, but should be affordable, give results rapidly, and be comfortable enough for repeated use.

Recommendation: Establish a rapidly responding contact-tracing and quarantine program. Even a small program could have significant impact if it has enough capacity to address those suspected of being super-spreaders (e.g., the top 5% of cases).

7 Is the campus an island from the community?

SARS-CoV-2 can spread in any human-to-human interaction across a given day. University personnel generally interact with community members at home, at the store, on public transportation, etc. Thus, infection rates for on-campus personnel are not exclusively a function of on-campus activities. This section explores how infection rates on a university campus are influenced by infections off campus.

7.1 Island model

We consider whether reducing the basic reproduction number R , within a small sub-population is effective in altering the course of the disease when that sub-population is embedded within a host population where the infection is less-well controlled. In an SIR epidemiological model (Appendix B), if we label the smaller university population as group 1 and the larger population as group 2, this interaction can be represented for group 1 with an effective transmissivity (β) given by

$$\beta^* = \beta_1 d \frac{I_1}{N_1} + \beta_2 (1 - d) \frac{I_2}{N_2}. \quad (7-6)$$

The terms, I/N , represent the fractions of groups 1 and 2 that are infected, d is the fraction of time group 1 is a distinct entity not interacting with other populations, and $1 - d$ is the fraction of time that group 1 participates as part of group 2. All fractions range from zero to one.

Simulations in Appendix B.2 illustrate how the influence of a growing infection in group 2 on infections in group 1 can be minimized by decreasing interaction with group 2 and achieving a smaller value of R in group 1. Given at least some interactions with group 2, however, infections in group 1 will inevitably infect those in group 2 whether or not group 1 keeps $R < 1$. Nevertheless, equation 7-6 makes clear that reducing the interactions with the broader community is helpful. Certain universities may find it valuable to develop programs that help their populations reduce community interactions, especially for high-risk groups (e.g., immune-compromised or mission-critical staff) and those living in tight quarters with other community members (e.g., dorms and group houses) where an infection poses a higher

risk to others. Such programs might include shuttle-bus services, childcare, provisions for groceries, lunchtime food, or access to off-campus business and services through programs that reduce risky interactions.

Disease dynamics in any population are, of course, the aggregate of the dynamics within the sub-populations. Efforts to control spread within a university will have benefits for the larger population so long as policies do not adversely affect the larger community. It follows, for example, that an infected individual that is detected by a screening activity at the university should quarantine, as opposed to continuing to interact with the larger population.

7.2 Archipelago model

The concept of isolating, at least partially, a sub-population in order to better control spread of the disease can be extended to individual groups within a university. Karin et al. [37] advocate a strategy involving dividing a university population into two halves, where one half works for 4 days in a given week followed by 10 days of quarantine. The other half works for 4 days in the alternate week, also followed by 10 days of quarantine. The advantage of this approach is that it lowers population density on campus and, even if the campus environment does not achieve $R < 1$, can help maintain an effective $R < 1$ through the quarantine phase.

The dynamics of the “4-10” strategy are analogous to those given by equation 7-6 if group 2 is re-interpreted as quarantined. In this case, $d = 4/10$ and β_2 is small. If, however, the “quarantine phase” becomes tantamount to normal interactions with the larger population, then the 4-10 strategy is helpful for reducing on-campus population density, and thereby reducing β_1 , but increases interactions with the larger population, making the university more beholden to the dynamics of the larger community. A more complete discussion of the dynamics of the 4-10 strategy is given in Appendix C.5.

The foregoing scenarios are simplistic representations of populations dynamics that involve complex spatio-temporal patterns and stochastic interactions. Populations and sub-populations are not thoroughly mixed, for example, and geometric aspects such as the relative location of a university and its host community will vary widely. Nevertheless, modeling illustrates scenarios of interest and helps show connections in the dynamics common to various policies. No university is an island, but efforts to control the spread

of the disease within such a sub-population have direct benefits.

Findings and recommendations

Finding: A standard epidemiological compartment model modified to represent a sub-population hosted within a larger population indicates that the fraction of infected in the sub-population is lower inasmuch as the sub-population both reduces contact with the larger population and maintains a lower R than the larger population.

Finding: Because disease dynamics in models have exponential growth rates, making a decision to move to a lockdown phase earlier, even by one or two days, can have significant positive impacts, reducing the duration of lockdown required.

Finding: The recently discussed “4-10” scenario for return-to-work is effective in lowering on-campus population density but its overall efficacy depends on persons in the 10-day quarantine period significantly minimizing interactions with non-university communities having a higher R .

Finding: Universities will influence and be subject to the disease dynamics of the larger communities in which they are embedded.

Recommendation: Universities should seek to understand the exposure of their personnel and the infection rates both on and off campus. Frequent coordination with local health officials will keep university decision-makers informed.

Recommendation: Universities may wish to develop programs to help people reduce their dependence on community engagements, especially those activities that pose a higher risk of infection, and for those individuals that are in vulnerable populations.

8 What operational policies are recommended?

Universities are developing operational policies allowing research to restart while attempting to manage risk to their researchers and surrounding community. The operational policies will require unfamiliar behavior on the part of researchers and the university administration will have the challenge of enforcing these new rules.

8.1 Communication

Informative communication is a key component of success. Researchers subject to new rules and procedures must understand them if they are to buy into their legitimacy. Training videos, written explanations, clear signage, and conversations with laboratory leaders should all be employed to educate the research community and enlist their cooperation.

Recommendation: Administrators, facilities, and health professionals should work to create materials to educate the research community on the new rules and procedures and they were devised.

8.2 Basic source of risk

Infection occurs when virus comes into contact with mucus membranes such as the eyes, nose, mouth or respiratory tract. The virus can be carried on objects (fomites), or in liquid droplets and aerosols suspended in the air. These virus-bearing substances impinge on a person as a series of individual events, each one bringing some probability of infection. The exposure rate $n(t)$ is the number of virions (virus particles) impinging on a person per unit time. The total dose is then

$$N_{tot} = \int_0^T n(t) dt.$$

The goal of risk mitigation measures is to reduce $n(t)$ and T to levels where the disease is unlikely to spread on campus, defined as $R < 1$ (see Section 1).

Finding: Dose = exposure rate $\times$ exposure time. Minimizing dose means minimizing the exposure rate and exposure time in all work environments,

including transit areas such as corridors, stairwells, and elevators.

8.3 Organizational methods to reduce dose

The findings above suggest simple and effective methods for reducing exposure in case an infected researcher makes their way into a research lab. Minimizing the time in the lab and working in shifts both provide means of reducing exposure. Creating research cohorts or “pods” may be possible for larger labs. Members of one pod should not mix with members of other pods, confining any potential infection (see Section 7 and associated appendices).

Each lab will need to find its own means of operating during a restart. While distancing, masks, and disinfection measures help reduce the risk, options that minimize the time two or more researchers are in a lab together will likely provide the most effective protection. Core or shared facilities, and remote labs such as telescopes or seismic stations, can all be made safer using the same principles applied appropriately at each location.

Recommendation: Adopt organizational changes to create shifts and cohorts that minimize the number of researcher interactions. Tailor these protocols to the needs of each laboratory, shared facility, and remote station.

Recommendation: Track which coworkers work together to facilitate contact tracing.

Recommendation: Develop systems and protocols that can minimize time spent in lab.

8.4 Respiratory masks

The primary viral transmission vector for COVID-19 is believed to be airborne transmission from droplets and aerosols (Section 3 and reference [14]). Large droplets (diameter greater than approximately $10\text{ }\mu\text{m}$) are actively generated by speaking, sneezing, and coughing. Droplets larger than $30\text{--}40\text{ }\mu\text{m}$ (at creation) are large enough to fall under the force of gravity and the probability of transmission is therefore a decreasing function of distance from the infected person [67]. Aerosols are smaller airborne particles produced by

breathing, and in much greater numbers by speaking. Once generated they can remain airborne for several hours, potentially causing transmission over larger distances and, speculatively, through unfiltered air-handling systems.

In the United States, public-health advisories about the use of masks covering the mouth and nose to mitigate COVID-19 transmission evolved from the initial statements that masks were not useful, to an acknowledgment that masks could reduce the risk of transmission, especially from the wearer to others. Unfortunately, the benefits of masks in reducing infection risk for the wearer have not been communicated nearly as well. The scientific literature on mask use related to viral transmission [58], as well as recent work specifically related to COVID-19, clearly indicate that masks also significantly reduce the risk of transmission in both directions, including for a mask wearer in the vicinity of an infected person [16, 69]. The level of protection varies with the material, design, and fit of the mask. Thus, university restart plans should include an evaluation of what types of masks are acceptable and best suited for the scenarios under consideration.

There are well-established standards for medical masks, which address the reduced transmission to the wearer from both aerosols and droplets, as indicated in table 1. Surgical masks are relatively comfortable and provide a useful benchmark for what may be achieved in masks to be made available to workers in most research environments.

If there are shortages, certified medical masks must be prioritized for healthcare workers. Other types of masks, both commercial and do-it-yourself (DIY), can be nearly as effective as surgical-masks in actual use, as described below.

8.4.1 Mask technical standards

The two key variables in masks are the materials used in the mask and the structural design of the mask itself. Public-health advisories for COVID-19 have recommended the use of cotton, which was the standard material used medically before the advent of disposable masks. Today surgical masks (rectangular-with-pleats) are made from special paper-based cloth, and most respirator style (fitted) masks are made from layers of electrostatically charged non-woven polypropylene fabric.

Some early literature on materials for DIY face masks did not provide

Mask Type	Standards	Filtration Effectiveness		
Single-Use Face Mask 	China: YY/T0969	Open-Data Tests Smart Air SmartAirFilters.com 3.0 Microns: ≥95% 0.1 Microns: ✗		
Surgical Mask 	China: YY 0469	3.0 Microns: ≥95% 0.1 Microns: ≥30%		
	USA: ASTM F2100	Level 1	Level 2	Level 3
		3.0 Microns: ≥95% 0.1 Microns: ≥95%	3.0 Microns: ≥98% 0.1 Microns: ≥98%	3.0 Microns: ≥98% 0.1 Microns: ≥98%
	Europe: EN 14683	Type I	Type II	Type III
		3.0 Microns: ≥95% 0.1 Microns: ✗	3.0 Microns: ≥98% 0.1 Microns: ✗	3.0 Microns: ≥98% 0.1 Microns: ✗
Respirator Mask 	USA: NIOSH (42 CFR 84)	N95 / KN95	N99 / KN99	N100 / KN100
	China: GB2626	0.3 Microns: ≥95%	0.3 Microns: ≥99%	0.3 Microns: ≥99.97%
	Europe: EN 149:2001	FFP1	FFP2	FFP3
		0.3 Microns: ≥80%	0.3 Microns: ≥94%	0.3 Microns: 99%

3.0 Microns: Bacteria Filtration Efficiency standard (BFE).

0.1 Microns: Particle Filtration Efficiency standard (PFE).

0.3 Microns: Used to represent the most-penetrating particle size (MPPS), which is the most difficult size particle to capture.

✗: No requirements.

Table 1: Compilation of technical standards for various types of medical masks. Source [54].

adequate statistics or material specifications to form generalized results. Several recent studies provide guidance on which materials are suitable [45, 51] provides some insight on the performance on material and designs. This study specifically addressed masks as-worn, which includes leakage around the edges of loose-fitting masks. The study used a 3M model 1826 surgical mask as a baseline for comparison. It demonstrated that for particles with a nominal size of 0.04 micron, the surgical mask removed more than 70% of particles, while rectangular single-use medical masks (Table 1) and rectangular cotton do-it-yourself (DIY) mask with a non-woven polyethylene insert removed almost 60%, and fitted cotton DIY masks removed 65–70%. Other DIY masks (rectangular cotton masks) tested more poorly, removing as little as 30% of small particles. This indicates the importance of specifying standards for masks.

Another recent study has illustrated the importance of how a mask is worn to its performance [69]. In this study, neither N95 masks nor surgical masks performed as indicated by their standards unless they were fit very tightly. The authors found that only by duct-taping the masks to the silicon dummy could the expected filtration efficiencies be achieved. Addition of a nose-clip to the surgical mask provided a significant increase in filtration efficiency. Masks must be worn properly throughout the working day if they are to provide the expected benefits.

If masks can be made available in bulk and certified for performance with proper fitting procedures, single-use medical masks could present a level of protection approaching that of surgical masks. However, these are disposable that may raise concerns for sustainability as well as continuing costs. If these items are not reliably available in the quantities needed, or if sustainability is highly valued, there are mask design and materials combinations that provide equivalent protection in washable (and thus reusable) masks. Providing a source of such masks would require some effort to establish specifications and identify a supply chain with sufficient capacity.

Finding: Mask efficacy is highly variable, and depends on the materials, designs, and user discipline in wearing the mask.

Recommendation: Universities should provide masks that meet demonstrated technical performance at the levels needed for the research environment. Training should be provided on how to properly wear masks.

8.5 Eye protection

To the degree that masks are advocated because they protect the wearer from infection (as well as protecting others from potentially being infected by the wearer), one can consider eye protection in addition to the mouth and nose protection offered by a mask. Safety glasses are needed for many laboratory applications, and can also help to protect the wearer's eyes from virus-bearing droplets in the air. Eyeglasses offer less protection than more encompassing safety glasses.

Face masks for respiratory protection can cause eye protection to fog up or become stifling. Since the evidence suggests respiratory protection is far more important, administrators might wish to consider the potential compliance problems generated by requiring eye protection in addition to respiratory protection.

8.6 Physical distancing and its limits

Some universities have created per-person area allocations in labs to enforce 6-foot physical distancing. While area allocations help reduce transmission via droplets, they are significantly less effective for aerosols that float in the air, presenting a risk to everyone in the room.

Section 3 contains a discussion of virus transmission by aerosols, and transmission mitigation by the HVAC systems and physical distancing. In the simplest terms, the room is a slowly leaking box. Limiting the time during which two or more people are in the room is the best way to prevent aerosol-borne transmission. In Section 8.4 we recommend that researchers always wear at least surgical masks, reducing the viral density in the room air, as well as the number of virus bearing particles inhaled. These measures make physical separation less critical, though still sensible whenever possible.

Finding: Standardized area allocations, minimum occupancy, and the use of masks work together to reduce dose by reducing exposure rate and exposure time.

Recommendation: Laboratory directors should work to minimize the occupancy of their labs to time when both people are essential to the task at hand.

Recommendation: Laboratory directors should consult their university's facilities and Health and Safety group on airflow in their labs to ensure there is at least adequate airflow and, and consider increasing the flow rate in the lab in case of higher occupancy.

Finding: Minimizing dose within the lab is insufficient; similar standards need to be met throughout the rest of the building, with particular attention to confined spaces that may receive little airflow but have frequent visitors (corridors, stairwells, and elevators).

8.7 Compliance

Universities will develop special rules for restarting research such as those recommended in this report. The restart rules add to the burden of usual rules from OSHA, NIOSH, the state, the city, and the university that govern laboratory operations. Laboratory directors must make it clear to their researchers that the restart rules are in addition to the usual laboratory rules, and that the restart rules may preclude activities normally allowed by the university. Laboratory directors may not have prior experience enforcing laboratory rules. Lab-level enforcement may become more important, for example, in ensuring continued face mask compliance. University administrations should provide training and support to help laboratory directors enforce rules by providing training materials to inform subordinate researchers, as recommended above.

Administrators must work with faculty and laboratory directors to develop meaningful consequences for those that do not follow rules and procedures. The consequences for misbehavior should be like any other safety violation and lead to expulsion from the laboratory building after warnings. As with the rationale for the creation of COVID-19 rules and procedures, the consequences and the reasons for these consequences must also be communicated through a variety of channels.

Leaving the enforcement of new rules and procedures to principal investigators (who may not even be on campus) or relying on reporting by other researchers can be expected to result in wide-spread flouting of the rules and procedures. Administrators must work with faculty, lab directors, the university's general counsel, and human resources to create a tiered response to rule flouting. Ultimately, the university should seek to create the sense that flouting is a transgression against the community, and sanctions must

be seen as coming from the university and not just from one's immediate supervisor.

Recommendation: Create and communicate a clear set of consequences for failure to follow COVID-19 rules and procedures and create a tiered response for transgressions.

9 How can institutions make risk-informed decisions?

Once a restart effort begins, institutional administrators will have to make regular decisions about expanding or contracting operations based on data gathered during the restart. This section considers how principal investigators and administrators may use this data for daily or weekly decision making about the conduct of operations.

Universities will operate in the low prevalence, low-testing regimes during the early stage. As such, insights into how campus operations are directly influencing the transmission of the disease will be hard to come by, as explained in Section 4.2.3. Instead, other data collected during a restart, such as information on how many people are in rooms, and how long researchers are in their labs, may be more actionable. Daily symptom attestation (see Section 5) can also be aggregated and studied for hints of off-campus infection.

Compliance data should also be collected to determine whether researchers are following new health-and-safety rules. Levels of adherence to assigned work hours and work areas, the proper wearing of masks, and consistency of symptom attestation can be thought of as leading indicators of the infection rate. In other words, not following the most basic rules should be expected to lead to increased levels of infection [25].

Since viral testing may be prohibitively expensive or of limited value because of false positives, the aggregate of other data, suitably processed and presented in summary form as a “dashboard” may help decision makers reduce the rate of infections on campus and know when to ramp up (or down) the number of people on campus.

There are multiple benefits of this proposed dashboard:

- Analysis of daily reports of health attestations, and whether or not those reporting symptoms subsequently test positive, may help refine algorithms and the associated weights given to symptoms.
- With time, improved virus tests with greater specificity and faster turnaround times will become available and results from such tests should be added to the dashboard and correlated with other data.

- Knowing who is on campus and where they are during the day can also inform contact tracing efforts
- Knowing which rooms were used, by how many, and when they are empty, can support cleaning and maintenance staff, allowing them to work safely.
- Further aggregation of information into a well-crafted community dashboard, accessible to all members of the institution’s community, is an opportunity to develop shared situational awareness of “how we are doing.”

Should there be evidence of increasing infection risk, recent studies illustrate the importance of rapidly locking down activities to prevent further disease spread [20, 30]. This value is also apparent from the history of the 1918 flu epidemic (Appendix A). It is critical that administrators and designated health officers have rapid access to the necessary data and monitor that data regularly, to identify a local outbreak so that they can make the needed decisions.

Findings and recommendations

Finding: Restart information related to symptom reporting, testing, campus access, and compliance with new rules must be rapidly reported and aggregated into a format that can support immediate decision making.

Recommendation: Create a public dashboard of testing, research compliance, facility access, symptom reporting, and local population prevalence information to inform decisions and create a shared sense of the situation.

A Lessons from the 1918 pandemic

It is not too much to ask what help we can get today from the experiences of our parents or grandparents in the H1N1 influenza pandemic of 1918–19. This pandemic, which killed tens of millions of people globally and an estimated 675,000 across the United States ⁶, stimulated a variety of nonpharmaceutical responses across our country. Here we summarize some lessons and warnings regarding these measures employed a hundred years ago in a variety of US cities.

Response to the 1918–19 pandemic in the US has received significant study, revealing some useful lessons and warnings that are relevant today. Although Covid-19 is not influenza, it is similar enough in transmission to be instructive regarding nonpharmaceutical intervention (NPI) – mainly social distancing measures. Similarly US cities reflect many geographic, social, economic and climate differences. Yet responses across many (but not all) cities were similar in the face of no effective anti-viral drug or vaccine – sound familiar?

A very useful statistical study of nonpharmaceutical intervention (NPI) methods, their effectiveness and application in mitigating excess death rates (EDR) was done using data culled from many sources of the time (1918–19) [46]. Their major findings from 43 cities demonstrate that early application; multiple techniques and sustained application are associated with mitigation of the EDR. Typical interventions were school closures, public gathering bans and isolation and quarantine. A number of studies indicate that these measures impacted time to peak death rate, peak death rate and total number of excess deaths [27]. In particular cities that implemented NPI earlier, used multiple interventions and sustained them longer reaped the benefits of delaying the peak EDR, reducing the peak EDR and reducing aggregate excess deaths. The conclusions of [46] are based on Spearman rank correlation studies: cities that implemented NPI earlier had greater delays to peak EDR (Spearman $r = -0.74$, $P < 0.001$), lower peak EDR (Spearman $r = 0.31$, $P = 0.2$) and lower total mortality (Spearman $r = 0.37$, $P = 0.02$). Further, there was a significant association between increased duration of NPI and reduced total mortality (Spearman $r = -0.39$, $P = 0.003$). In addition to the effectiveness of NPI the work of Markel et al. illustrates the dangers of relaxing NPI measures too soon and suffering a second wave of deaths.

⁶<https://www.cdc.gov/flu/pandemic-resources/1918-pandemic-h1n1.htm>

Fig. 7 illustrates a “tale of four cities” and how they fared with respect to start up date, types, duration and relaxation of nonpharmaceutical interventions. Probably the most interesting part of this tale is St. Louis, panel A in Figure 7. In early October, city health commissioner Dr. Max C. Starkloff ordered closure of schools, movie theaters, saloons, sporting events and other public gathering spots as well as suspension of Sunday church services [6]. In many respects St. Louis fared relatively well with the 8th lowest aggregate excess deaths per capita of 43 cities studied. However, this city also points out that relaxation of NPI measures too soon can lead to a resurgent, second wave. For St. Louis premature relaxation resulted in a second peak with death rate higher than the first peak. Denver in panel C illustrates a similar course of events with an early relaxation and subsequent larger second wave. In both cases reimposition of NPI measures was used and resulted in a steep decline in the death rate. New York City (panel B) fared relatively well with an early and sustained application of isolation and quarantine and an aggregate death total ranked 15th lowest of the 43 cities studied by [46]. Pittsburgh (panel D) had the highest aggregate death rate of all the 43 cities studied. The use of multiple NPI measures was limited and not sustained. Neither New York nor Pittsburgh suffered a pronounced second wave. In summary although there are many confounding factors of geography, climate, air pollution, etc., the association of NPI measures, now called social distancing, with reduced rates of respiratory virus disease transmission has been convincingly demonstrated in studies of the 1918–19 influenza pandemic. Ref. [46] argue that in future pandemics NPI “might play a salubrious role in delaying the temporal effect of a pandemic, reducing the overall and peak attack rate; and reducing the number of cumulative deaths.” This view is supported by a variety of other sources (e.g., [21] and [27].)

However, the lessons of how to apply nonpharmaceutical techniques and the pitfalls of ceasing these measures too early appear not to have been as widely recognized as historical lessons demonstrate in our “tale of four cities,” above. Further, we note that opponents of NPI techniques can have serious impact and disastrous effects. San Francisco in 1918–19 is an example where NPI was successfully introduced under a WW1 patriotic banner, but tended to wane as residence tired of restrictions and masks. Eventually after much controversy NPI measures were relaxed and a strong second wave occurred, bringing reports that San Francisco’s cumulative death toll was the highest of any major US city, estimated at 673 per 100,000.⁷

⁷“San Francisco, California and the 1918–1919 Influenza Epidemic.” University of Michigan Center for the History of Medicine: Influenza Encyclopedia, <https://www.influenzaarchive.org/cities/city-sanfrancisco.html>

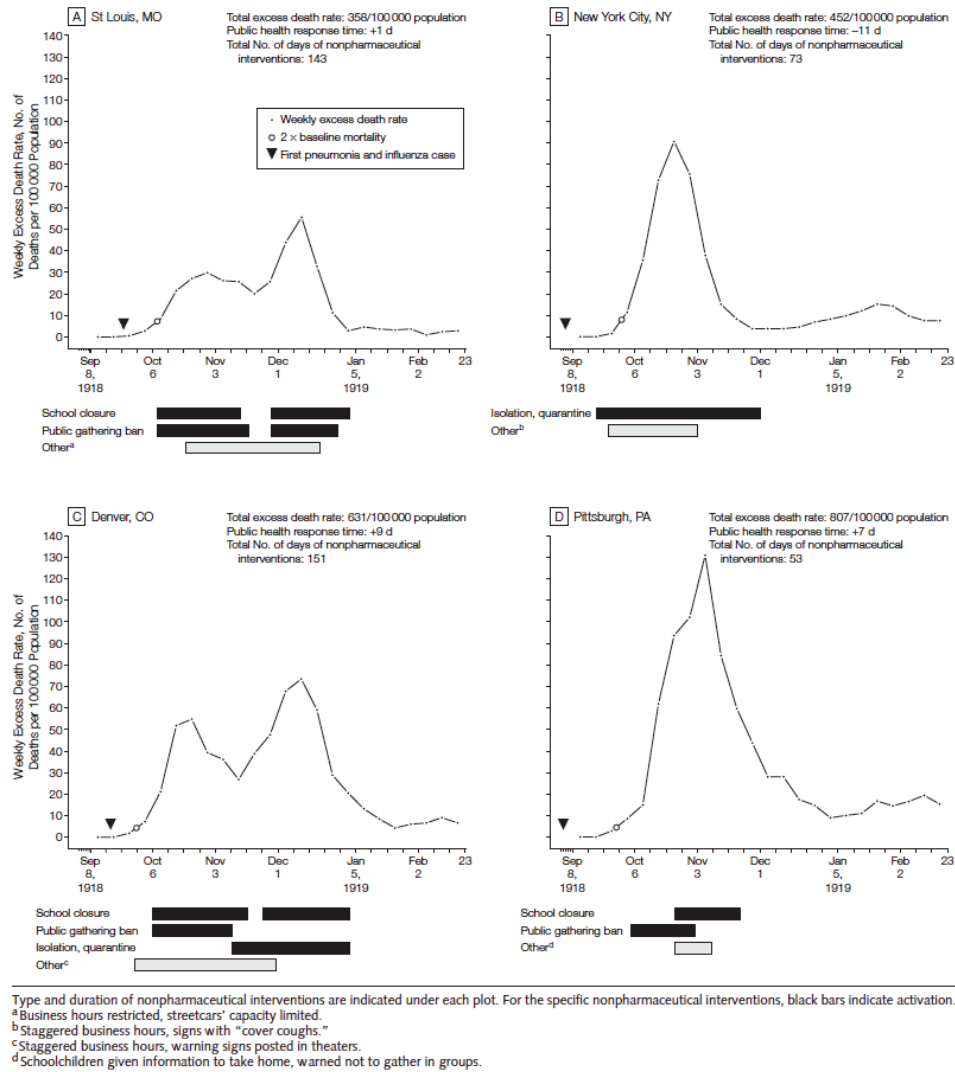


Figure 7: A tale of four cities—response to 1918–19 pandemic. Weekly excess death rates from Sept. 8, 1918 through Feb. 22, 1919. For these four cities the profile of weekly death rates are shown with types of NPI techniques employed and dates in use. Note the positive effect of early implementation of multiple techniques and sustained use and the negative effects of premature relaxation of NPI use and the resurgence in a second wave. After [46].

B Strategies to reduce R

An SIR model represents the population as residing within three compartments: susceptible (S), infected (I) and recovered (R). The sum of the fractional populations across these compartments is $S + I + R = 1$. The rate at which susceptible individuals become infected depends on the transmissivity, β , as well as the number already infected, and the rate at which infected recover, γ , giving:

$$\dot{S} = -\beta SI, \quad (\text{B-7})$$

$$\dot{I} = \beta SI - \gamma I, \quad (\text{B-8})$$

$$\dot{R} = \gamma I. \quad (\text{B-9})$$

In the approximation of a disease-free limit, $S = 1$, the solution to Eq B-8 is, $I(t) = I_o \exp[(\beta - \gamma)t]$, where I_o is an initial fraction infected. For I to grow, β must be greater than γ . Equivalently, the basic reproductive number, defined as $R_o = \beta/\gamma$, must be greater than one. R_o indicates the number of people an infected person subsequently infects under the idealized circumstances that the entire population is susceptible. Note that setting Eq. B-8 to zero permits for solving for S at peak infectivity, which equals $1/R_o$.

Within the context of the SIR model there are only two general approaches for reducing R_o : increasing the rate of removal and decreasing transmissivity. First, we consider increased rates of ‘removing’ infected via quarantine in the context of SIR. Second, we consider whether a subgroup that decreases β , or increases γ , can control the dynamics of the disease when embedded within a larger population, again using SIR. These simulations quantitatively illustrate how quarantine and reductions in transmissivity can be effective for controlling the spread of COVID-19.

B.1 Increasing the rate of removal

We wish to update Eq. B-8 to include the effects of identifying infected individuals through symptom screening or testing and subsequently quarantining those individuals such that they do not interact with the susceptible population.

B.1.1 Screening via symptom attestation

We assume a rate of symptom attestation of r_s and that this approach to screening has a false negative rate of f_s . Assuming that people identified with symptoms are quarantined, we have,

$$\dot{I} = \beta SI - \gamma I - r_s(1 - f_s)I. \quad (\text{B-10})$$

The solution in the disease free limit is $I(t) = I_o \exp[(\beta - \gamma - r_s(1 - f_s))t]$, such that the analogue of R_o becomes $\beta/(\gamma + r_s(1 - f_s))$.

A baseline SIR simulation is adopted for illustrative purposes having an $R_o = 2$, with $\beta = 1/(3 \text{ days})$ and $\gamma = 1/(6 \text{ days})$ (Eq. B-8). The baseline simulation gives peak infection rates of 15% and 80% of all people ultimately become infected. If Eq. B-10 is instead applied with a symptom attestation - based screening rate $r_s = 1/(2 \text{ days})$ and $f_s = 80\%$, peak infections are 2%, 37% are ultimately infected, and $R_o = 1.25$ (see Fig. 8).

A larger value of r_s than γ reflects the potential for screening a population regularly, and the large value of $f_s = 80\%$ reflects imperfections of symptom attestation as a screening approach. To take screening according to fevers as an example, more than 80% of patients have been reported to present with fevers at the onset of the disease [44, 11] but there is a significant risk that asymptomatic individuals can also transmit the disease [24]. Errors in screening equipment or personnel performance as well as ignoring or evading screening measures could further contribute to false negatives. Gostic et al. [26] estimate that even under best-case assumptions that screening will miss more than half of infected people.

There are two other factors that inform our use of a high false negative rate. First, the false negative rate will evolve over the age of an infection. How early an individual can be detected is critical for stemming the spread of the disease, such that individuals that are infectious but asymptomatic pose a major challenge [24]. The simple SIR model used here, however, treats the probability of identifying anyone that is infected as being equal. Second, and related to the first point, our representation assumes that each screening is independent, with an increased rate of screening proportionately increasing rates of removal. In fact, repeated screening of an infected individual will only allow for quarantine once symptoms appear. A more complete screening model would account for emergence of symptoms together with the rates of screening when determining rate of removal. Adjusting the model to rates faster than $1/(2 \text{ days})$ would further strain the assumption that tests are inde-

pendent. A final consideration is that screening will typically be undertaken using a more complete set of symptoms than just fever, such as including cough, shortness of breath, or fatigue. The false negative rate when using multiple symptoms for screening is lower relative to using just one [47], but the degree to which false negative rates are suppressed will depend upon the degree to which symptoms are independent from one another and the specific criteria used when calling for quarantine.

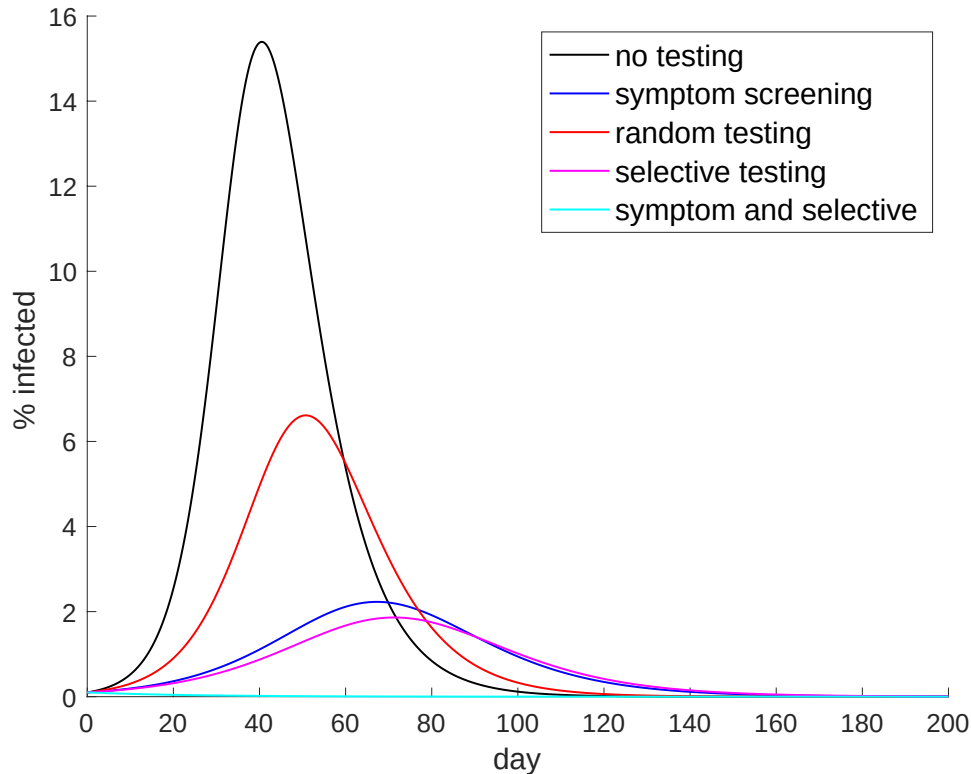


Figure 8: Infections from a baseline SIR model (black) and with introducing symptom attestation (blue), random testing (red), focused testing (magenta), and symptom attestation and focused testing combined (cyan). Baseline values have $\beta = 1/(3 \text{ days})$ and $\gamma = 1/(6 \text{ days})$ with all modification acting to increase the rate of removal and, hence, the effective γ .

B.1.2 Testing

From the perspective of reducing R_o , testing is complimentary to symptom attestation, and can be represented similarly,

$$\dot{I} = \beta SI - \gamma I - r_t(1 - f_t)I. \quad (\text{B-11})$$

Assuming that the population can be tested at a rate of $r_t = 1/(14 \text{ days})$ and $f_t = 25\%$ false negatives gives peak infections of 7%, 59% ultimately infected, and $R_o = 1.5$. For the selected parameters, screening the population for symptoms every 2 days appears more effective than testing every 14 days despite testing having a much lower false negative rate. The assumptions of independence and equal probability associated with our simple screening model is also made with respect to testing. Confidence in the conclusion that screening would out-perform testing is tempered by these assumptions. Specifically, independence of test results when applied at a rate of $1/(14 \text{ days})$ seems relatively more plausible than independence of screening at a rate of $1/(2 \text{ days})$.

Eq. B-11 assumes randomized testing with the result that the number of tests administered to infected individuals is linearly proportional to I . In practice, a more common testing strategy than randomized testing is to administer tests to individuals that are more likely to be infected. For example, tests could be preferentially administered to those whose work or travel requirements make them more likely to be exposed or those who have been in contact with symptomatic individuals. Antibody testing could also be used to exclude those who recovered from infection.

To account for selective testing, we define the probability that someone is infected conditional on being tested as,

$$P(I|T) = \frac{P(T|I)P(I)}{P(T)}, \quad (\text{B-12})$$

where the right hand side uses Bayes Theorem. The probability of testing is,

$$P(T) = P(T|I)P(I) + P(T|S)P(S) + P(T|R)P(R). \quad (\text{B-13})$$

Assuming that the conditional probabilities are constant leads to,

$$P(I|T) = \frac{I}{I + \frac{p_s}{p_i}S + \frac{p_r}{p_i}R}. \quad (\text{B-14})$$

Eq. B-14 represents the ability to concentrate testing on the infected population.

Modifying Eq. B-11 to account for selective testing gives,

$$\dot{I} = \beta SI - \gamma I - \frac{r_t(1 - f_t)I}{I + \frac{p_s}{p_i}S + \frac{p_r}{p_i}R}. \quad (\text{B-15})$$

Ratios $\frac{p_s}{p_i}$ and $\frac{p_r}{p_i}$ describes the degree to which tests are focused away from susceptible and recovered populations onto the infected. Randomized testing corresponds to p_s , p_r , and p_i being equal, where upon Eq. B-15 simplifies to Eq. B-11. Assuming that $\frac{p_s}{p_i} = \frac{p_r}{p_i} = 0.5$ doubles the efficacy of testing and gives $R_0 = 1.22$, slightly below that simulated for screening. A similar equations could be applied to symptom attestation, but the fact that this approach to screening is generally cheaper, faster, and less invasive than testing suggests that such focusing is less important to consider.

Combining both screening and testing according to the foregoing parameter specifications gives an $R_o = 0.89$, such that the initial small number of infections in the population decays. This combined results illustrates that a screening strategy that combines symptom attestation and testing will, generally, be more effective than either in isolation. Here, it is assumed that symptom attestation and testing act independently, but correlations in false negatives would yield smaller improvements.

B.1.3 Pulsed testing

In seeking an optimal approach to testing the question arises as to whether a time-variable testing regime would help further reduce the spread of the disease. Fig. 9 shows a simulation having the same parameters as for the selective testing simulations shown in Fig. 8 but also includes a scenario whereby testing maintains its long-term average but alternates between two weeks of intensive testing and two weeks of reduced testing. The growth of the disease accordingly alternates from declining during intervals of intensive testing to rapid growth during intervals of reduced testing, but no long-term change in the course of the disease is apparent. Both scenarios produce 22% of the population having recovered from infection after 1000 days.

Results are qualitatively unchanged in our simulations using other choices of β and γ , other amplitudes of the square wave, different periods of the intensive-relaxed testing procedure, or use of sinusoidal variations. Including an incubation compartment in the model, i.e., an SEIR model, slows the trajectory of disease spread but appears similarly insensitive to constant versus pulsed testing.

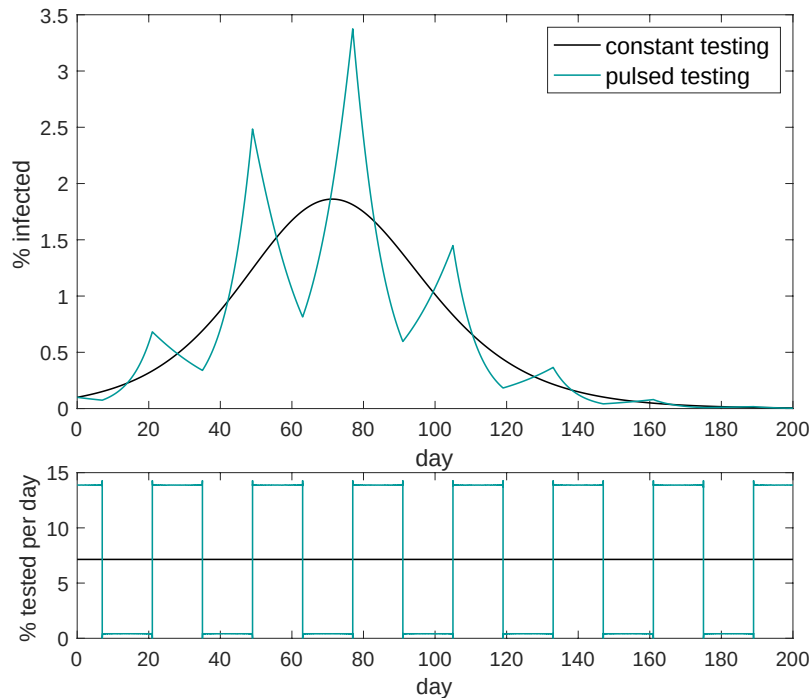


Figure 9: Simulation of the percent infected in a population using selective testing using a constant rate of testing and pulsed testing (blue). (Bottom) Constant and a pulsed testing scenarios; both have the same long-term average.

B.2 Sub-population

A variety of means exists by which to reduce transmission, which are considered in other sections, and to increase removal, some of which are illustrated in foregoing subsections. The question arises as to whether reducing R_o within a sub-population is effective in altering the course of the disease when it is embedded within a larger population with which the sub-population interacts.

We consider a model having a small sub-population, group 1, that interacts with a larger host population, group 2. Group 2 evolves independent of group 1, whereas group 1 partially follows its own dynamics for a fraction of the day, d , and otherwise that of the larger group, $1 - d$. Transmission becomes a weighted average between the β 's associated with each group,

$$\beta^* = \beta_1 d I_1 / N_1 + \beta_2 (1 - d) I_2 / N_2 \quad (\text{B-16})$$

. The transmission characteristics of group 1 become more important as the

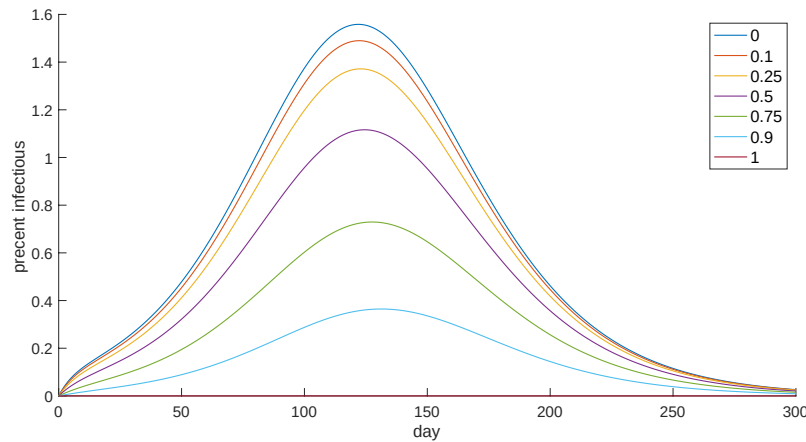


Figure 10: Simulation of the percent infected in group 1, where members of group 1 spend between all their time with group 1, $d = 1$, or all their time with group 2, $d = 0$. Insomuch as $d > 0$, members of group 1 continue to inherit infections from exposure to group 2.

fraction of time spent in group 1 increases and the prevalence of infection in group 1 grows relative to that in group 2. The SIR model for group 1 becomes,

$$\dot{S}_1 = -\beta^* S_1, \quad (\text{B-17})$$

$$\dot{I}_1 = \beta^* S_1 - \gamma I_1, \quad (\text{B-18})$$

$$\dot{R}_1 = \gamma I_1. \quad (\text{B-19})$$

Group 2 follow basic SIR dynamics irrespective of variations in Group 1. Simulations are made using $\gamma = 1/(6 \text{ days})$, $\beta_1 = 1/(10 \text{ days})$ and $\beta_2 = 1/(5 \text{ days})$. Infections would decline in group 1 except for ongoing exposure to group 2, which is parameterized with d ranging from 0 to 1 (Fig. 10).

Although the models explored here are very simple, they point to the importance of symptom screening and testing, followed by quarantine of positive cases, for purposes of mitigating the spread of disease. Focusing tests on those most likely to be infected is helpful, whereas time-variable deployment of testing capability shows no discernible advantage. Measures taken within a sub-population are effective insomuch as interaction with a much larger population do not dominate transmission characteristics.

C Shift work and community interactions

This appendix summarizes the predictions of standard epidemiological modeling using the standard disease basic reproduction number R as a measure of the spreading characteristics. R is a combination of the transmissibility of a disease along with the number of potential contacts, which can be reduced via social distancing measures. We include the basic population responses indicated by the time course of infectious individuals during the development of an epidemic followed by a lock-down with a much reduced R .

C.1 Spread of a virus

It is worthwhile to review a few facts about how a virus propagates in a population as these features are critical to both the spread of the virus, as well as a means to defeat the virus in the absence of a vaccine, e.g. [50]. Typically a virus reproduces in new hosts every τ days, producing R_0 new infections, where R_0 is referred to as the reproduction ratio or reproduction factor. To emphasize the rapid (exponential) growth possible, if $\tau = 3$ days and $R_0 = 3$, which are approximately the case for SARS-CoV-2, then 10 infected individuals on day one produce 30 infections by day four, and 90 infections by day seven (a factor of 9 in one week). After one month we can expect a multiplication of the virus by approximately $3^{10} \approx 59,000$ or, given the starting number of 10 infections, a total about 600,000 infected individuals. If the mortality rate is say one percent (data for SARS-Cov-2 is likely higher), than that is 6000 deaths in just one month from this one disease. The numbers also increase rapidly. It is clear that when such a virus enters a community, it is necessary to take aggressive action to impede the spread.

The reproduction ratio R_0 can be thought of as the product of the number of contacts a healthy, or susceptible, person has with an infected person times the probability of acquiring the virus; the latter almost certainly increases with the time of contact between individuals; for $R_0 > 1$ the virus spreads through the population, with an exponential growth in the number of infected individuals as the previous example indicates. Social distancing measures, with the extreme being a lockdown on a community, correspond to a time period during which the effective reproduction number has decreased significantly and reached values < 1 so that the number of new cases per day in the community is decreasing. Upon return to work we can expect the

effective R_0 to increase above its value from the time period of lockdown and if its value exceeds unity a second wave of the virus is expected. If the latter occurs, the sooner actions are taken to reduce R_0 , the better the chance of preventing the rapid growth of infections that happened in the first wave.

One important feature to recognize is that R_0 is, in part, a *social* number. It combines features of the virus, e.g. its infectivity, with how people in a community interact. Although during the spread of the virus, the estimates of the reproduction ratio are $R_0 \approx 2.4 - 3$, it is possible that social distancing strategies, e.g. masks, keeping some distance away from people during conversations, washing hands regularly, etc. can significantly reduce R_0 , though it is not easy to estimate this reduction of the individual steps or the aggregate of steps. It is clear that mitigation strategies will be more challenging in high population density regions than in low density communities.

We first illustrate the state of a lockdown to see how the effectiveness of a lockdown impedes the spread of a virus.

C.2 Basic Model: SEIR for a Single Group

We assume that the population has a fixed number of individuals and for convenience work with equations representing fractions of the population. We use a common epidemiological model, which is a system of first-order ordinary differential equations for the parameters S , E , I , and R , where S indicates the fraction of the population that is susceptible to the disease, E is the fraction of the population exposed but as yet not infectious, and I is the fraction of the population who are infectious but not yet symptomatic. R denotes the fraction of the population who are either symptomatic but not circulating to infect others (i.e., they are confined to their homes), have recovered, or have succumbed to the illness. The equations describing the

system are:

$$\dot{S} = -\frac{1}{\tau_I} R_0 S I \quad (\text{C-20a})$$

$$\dot{E} = \frac{1}{\tau_I} R_0 S I - \frac{1}{\tau_E} E \quad (\text{C-20b})$$

$$\dot{I} = \frac{1}{\tau_E} E - \frac{1}{\tau_I} I \quad (\text{C-20c})$$

$$\dot{R} = \frac{1}{\tau_I} I. \quad (\text{C-20d})$$

Here R_0 describes how strongly S and I interact in transmitting the disease, τ_E is the characteristic amount of time a person remains exposed before becoming infectious, and τ_I is the typical time a person remains infectious before showing symptoms. Following Karin et al.[37], in the simulations reported here, $\tau_E = 3$ days and $\tau_I = 4$ days, as suggested by data reporting COVID-19 transmission. It is assumed that after a person show symptoms and transfers from I to R , they become sufficiently well isolated that it is acceptable to approximate them as no longer interacting with the broader population. The population fractions $S(t)$, $E(t)$, $I(t)$, and $R(t)$ are normalized so that $S(t) + E(t) + I(t) + R(t) = 1$. For example, $I(t)$ always denotes the fraction of the population at any time t that is infectious. Similarly, at any time t the fraction of the population what will soon be, is or has been infected is $E(t) + I(t) + R(t) = 1 - S(t)$.

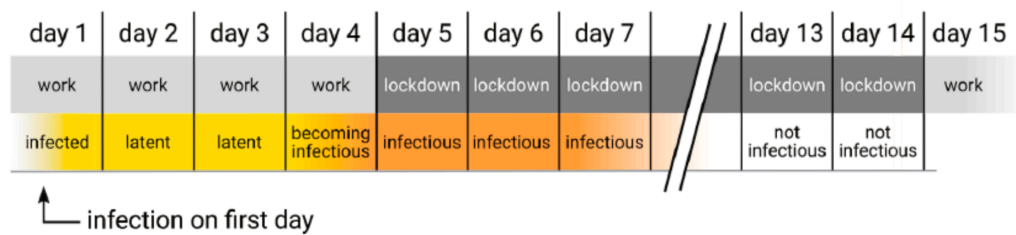


Figure 11: Evolution of an infection in SARS-CoV-2. Reference: [37]

C.2.1 Time Course for Development of Symptoms of COVID-19

We utilize a SEIR model (susceptible-exposed-infected-removed), where the fraction of the population in the different groups is represented by $S(t)$, $E(t)$, $I(t)$

and $R(t)$. Two time constants represent the features of a virus that infects a person. An individual remains asymptomatic for a few days, yet can spread the virus, and then only somewhat later shows symptoms of the virus. The typical time variation of infection characteristics of COVID-19 means an infected individual is asymptomatic but infectious 3 days after exposure and has a peak infectiousness occurring about four days after exposure (Figure 11). In the mathematical model the presence of symptoms is assumed to remove the worker from the work environment. The assumption are approximately consistent with the current understanding of the virus, e.g. Lauer et al. estimate a mean incubation period of 5 days, i.e., the time an infected person has from exposure to showing symptoms (with $< 3\%$ showing symptoms within 2 days and $> 97\%$ showing symptoms within 11 days [17]).

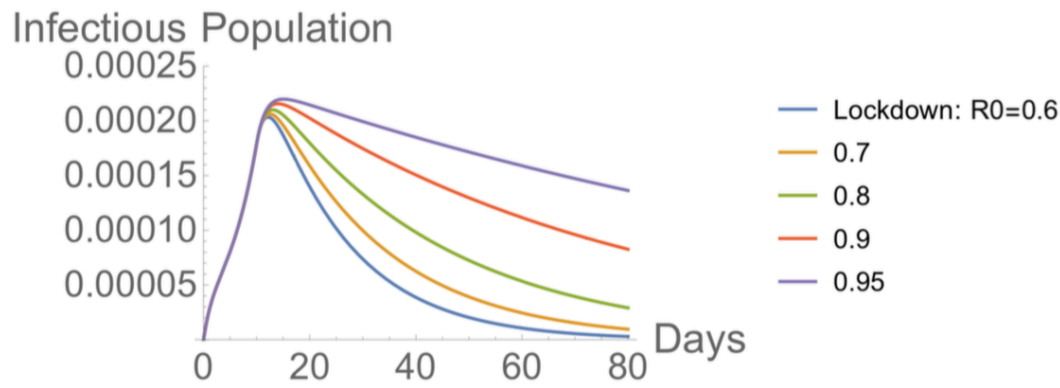


Figure 12: The infectious population $I(t)$ in the evolution of a virus, with a time of epidemic growth followed by a lockdown to slow the growth. In the simulation the epidemic evolved for 10 days with $R_0 = 2.4$, which was then followed by a lockdown with $R_0 = 0.6, 0.7, 0.8, 0.9$ and 0.95 , as shown by the different curves. The simulation utilized an SEIR model with times scales $\tau_E = 3$ days and $\tau_I = 4$ days and an initial value $S(0) = 0.9999$ and $E(0) = 0.0001$.

C.3 Growth of an Epidemic and Lockdown

To appreciate the typical dynamics leading to a return-to-work scenario, we use the well-known SEIR model to simulate spread of a virus with $R_0 = 2.4$. In some major cities a partial or complete lockdown was imposed one or two weeks after infections were recognized in the community, though the virus

was likely spreading in the community for a longer time. For example, in New Jersey, on 16 March there were approximately 80 recognized infections of SARS-CoV-2, and the Governor issued a stay-at-home order effective 21 March.

In the example simulations shown in this section, after 10 days of spreading of the virus, lockdown of the community is imposed and, for simplicity, though it does not affect qualitative features, R_0 was set to a constant value less than unity. We chose $R_0 = 0.6 - 0.95$ and report the fraction of the population that is infectious $I(t)$ as a function of time in Figure 12. We observe that $I(t)$ first grows rapidly (exponentially), peaks a few days after lockdown begins and then progressively decays. Note that for this example, the peak in the fraction of the population that is infected is 0.0002. In a community of $10M$, this corresponds to 2000 individuals who are infected circulating in the community on an given day.

Not surprisingly, following lockdown the rate of decay is tied to the effectiveness of the lockdown ($R_0 < 1$). It can be shown analytically in the simple SIR model that in the case the $S(t) \approx 1$ (only a small fraction of the population gets the virus), then for times after lockdown, $t > t_{LD}$, the infectious fraction changes according to $I(t)/I(t_{LD}) = e^{((R_0-1)(t-t_{LD})/\tau)}$. Use of published data from virus testing is then one way to estimate R_0 during a period of lockdown, at least if the number of positive cases can be reasonably associated with the fraction of the population that is infectious.

In addition, we report the total fraction of the population that will soon be, is or has been infected, $1 - S(t) = E(t) + I(t) + T(t)$, as shown in Figure 13. Some fraction of this group will have been hospitalized or will need hospitalization, of which a smaller fraction will end up in the ICU; a small fraction of the infected will die. We can see that an exponential, early time growth period transitions, upon lockdown, to a stage of nearly linear growth in the total number of infected individuals. For $R_0 < 1$ in lockdown we expect the virus to be defeated, at least temporarily, and the rate of diminution is larger for smaller R_0 . Note also that for a successful lockdown such as $R_0 = 0.7$ in Figure 13 the the fraction of the population that has been infected at 80 days reaches about 0.0015, which in a community of $10M$ corresponds to 15,000 people.

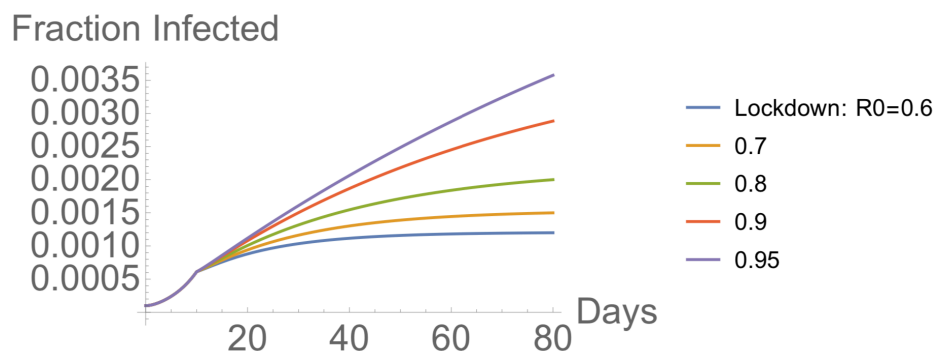


Figure 13: Total fraction of the population that will be, is, or has been infected at any time, $1 - S(t)$, obtained from the same simulations that led to Figure 11. In the simulation the epidemic evolved for 10 days with $R_0 = 2.4$, which was then followed by a lockdown with $R_0 = 0.6, 0.7, 0.8, 0.9$ and 0.95 , as shown by the different curves. The simulation utilized an SEIR model with times scales $\tau_E = 3$ days and $\tau_I = 4$ days and an initial value $S(0) = 0.9999$ and $E(0) = 0.0001$.

C.3.1 An earlier lockdown

In the face of a spreading virus rapid action is key. This will be important if a second wave were to start but the idea is illustrated with the dynamics at the start of the epidemic. We use the parameters of the previous example, $R_0 = 0.7$, but implement a lockdown after 5 days or 8 days, which can be compared with the lockdown after 10 days. The results for the total fraction infected (again, in the SEIR description, we calculate this as $1 - S(t)$) as a function of time are shown in Figure 14. The transition away from the exponential growth following lockdown is apparent. As in the previous example, if the parameters chosen for the simulation applied to a population of $10M$, then they predict 15,000 total infected individuals if the lockdown is effected after 10 days, but only about 6500 infected individuals if action were taken after 5 days.

C.4 Dynamics following a return to work

During successful lockdown the number of infectious individuals decreases. Any simple return to work strategy should be expected to increase the effective value of R_0 and when $R_0 > 1$ the virus will spread again since there

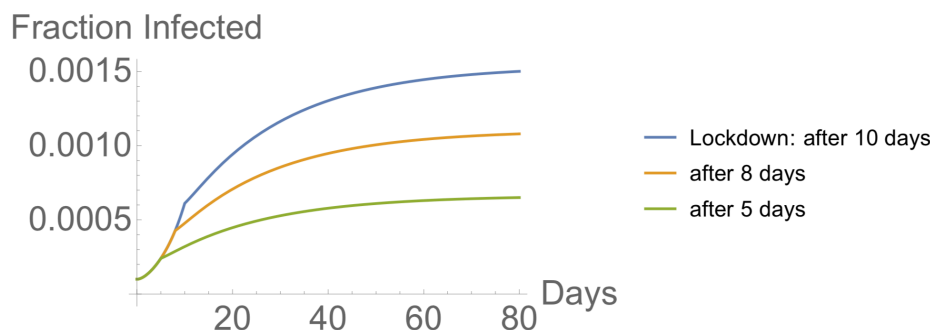


Figure 14: Fraction of the population infected versus time for lockdowns after 5, 8 or 10 days with $R_0 = 2.4$ initially, which was then followed by a lockdown with $R_0 = 0.7$. The simulation utilized an SEIR model with times scales $\tau_E = 3$ days and $\tau_I = 4$ days and an initial value $S(0) = 0.9999$ and $E(0) = 0.0001$.

remain infected individuals in the population to trigger a second wave. The rate of increase is faster the larger the value of R_0 is above unity.

To illustrate these points and to highlight how the magnitude of R_0 in the return-to-work phase can have significant influences on the rate of change of infections in the community we ran SEIR simulations with $R_0 = 2.4$ for 10 days, which was followed by a 40 day lockdown period with $R_0 = 0.6$ (see Figures 11 and 13 where the same initial conditions were used). We report how different values of R_0 in the return-to-work phase affect the time course of the fraction of the population that is infectious as a function of time (Figure 15).

For the values chosen for this simulation, at the peak of the initial phase, 0.02% of individuals were infectious (and present in the community) at about days 11-12, but by the end of the 40-day lockdown this number had decreased by a factor of 10. Nevertheless, even if $R_0 = 1.2$, which is estimated to be half of the current typical value for SARS-CoV-2 at the start of the pandemic, the number of infectious individuals has almost doubled 30 days later (day 80). By day 100, the increasing rate of infections is evident and the number of infectious individuals has already reached 1/3rd of the value at the peak near the beginning of the epidemic.

Only slightly larger values of R_0 lead to much larger growth rates, because these responses are exponential, with a rate approximately proportional to $R_0 - 1$. Thus, we can observe in Figure 15 that, when $R_0 = 1.4$ and 50 days

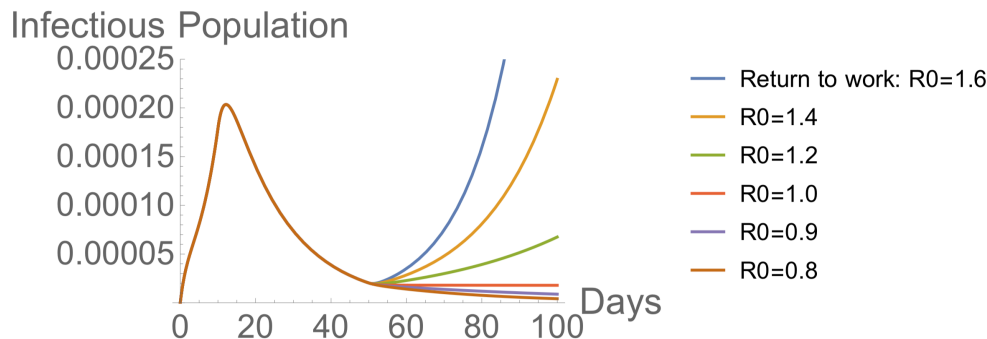


Figure 15: Fraction of the population that is infectious at any time following an epidemic, lockdown and a return-to-work phase. In the simulation the epidemic evolved for 10 days with $R_0 = 2.4$, which was then followed by a 40-day lockdown with $R_0 = 0.6$, after which a return-to-work period began. The simulation utilized an SEIR model with times scales $\tau_E = 3$ days and $\tau_I = 4$ days and an initial value $S(0) = 0.9999$ and $E(0) = 0.0001$.

into the return-to-work phase, already by day 100, the number of infectious individuals has exceeded the peak earlier in the epidemic.

Although these simulations are based on a simplistic model they do highlight that, in the absence of a vaccine, there is not a lot of room for error in encouraging, implementing, and/or enforcing all manners of social distancing strategies to try to maintain R_0 below unity in a return-to-work environment. Next, we discuss a strategy that was suggested recently to return to work using a time periodic work cycle tied to the evolution of the infection in an exposed individual.

C.5 Model of Karin et al.: A “4:10” Strategy,

C.5.1 The idea of the “4-10” work cycle

Because the virus has a dynamics that symptoms typically show up within 4-5 days (Figure 11), Karin et al. offered a strategy for people to work according to a cycle of k days in the office and $14 - k$ days in lockdown at home [37]. Before we show their full model we illustrate how this approach of returning to the office 4 of 10 workdays during a two week period (denoted “4:10”), with higher R_0 during those workdays, is compatible with maintaining the virus

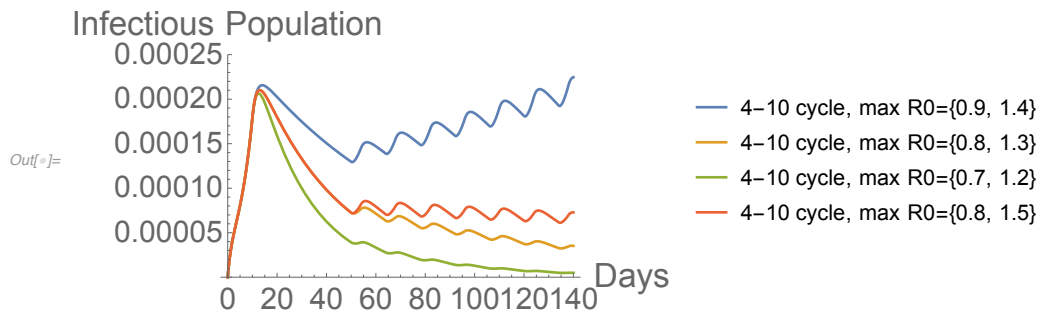


Figure 16: Infectious fraction for different variations of a “4:10” strategy in a back-to-work phase (individuals working 4 consecutive days every two weeks). In the simulation the epidemic evolved for 10 days with $R_0 = 2.4$, which was then followed by a 40-day lockdown after which a return-to-work period began. The values for R_0 during lockdown and workdays are indicated. The return-to-work phase includes time periods with $R_0 > 1$, yet the 4:10 strategy maintains an effective $R_0 < 1$ so that the infections continue to decline. The simulation utilized an SEIR model with times scales $\tau_E = 3$ days and $\tau_I = 4$ days and an initial value $S(0) = 0.9999$ and $E(0) = 0.0001$.

at bay, and not triggering exponential growth if the maximum R_0 , though greater than unity, is not too large. In the model the restricted number of work days has the feature that it maintains an effective value of $R_0 < 1$.

We ran the SEIR model as illustrated above (10 days of epidemic, then 40 days of lockdown, followed by a return to work), but this time implementing a return-to-work phase with with the “4:10” strategy. We choose various combinations of R_0 during lockdown and during workdays, respectively, as $\{R_0^{\text{lockdown}}, R_0^{\text{work}}\} = \{0.9, 1.4\}, \{0.8, 1.3\}, \{0.7, 1.2\}$. The state at the end of lockdown is shown in Figure 11.

The typical results for these values of R_0^{lockdown} are shown in Figure 16. In each case the workdays have $R_0 > 1$, yet for the two simulations with the smaller R_0 at lockdown, the fraction of infectious individuals, though now oscillatory, is continuing to decrease during the back-to-work phase. In the case of $\{R_0^{\text{lockdown}}, R_0^{\text{work}}\} = \{0.9, 1.4\}$ we see that the effective value of R_0 during the two-week cycle is sufficiently large that the fraction of infectious individuals increases.

This “4-10” strategy suggests that outside of lockdown social distancing should limit the incremental increase in the reproduction ratio to approxi-

mately $\Delta R_0 = \frac{7}{2}(1 - R_0^{\text{lockdown}})$. For $R_0^{\text{lockdown}} = 0.8$ this allows $R_0^{\text{work}} = 1.5$, for which the simulation is shown in Figure 16. In fact, the return-to-work curve is slowly decreasing in this case.

C.5.2 Two halves of the community in “4-10” cycles

One strategy towards a restart is to bring part of the work force back part of the time so as to maintain a work environment with a relatively low density, thus contributing towards reduced contacts. For example, the work force can be divided into two or three non-mixing groups and each group can work a set number of days followed by a period of workdays at home or a furlough. Karin et al. [37] used SEIR-type models to predict how the number of infected people (the workers) in a population varies in time if a work schedule is implemented with k days at work and $14 - k$ days of social isolation (“lockdown”), i.e., at the end of each day the workers return home during which they are assumed to have a lower level of contacts than in the workplace (e.g. family, shopping, etc.).

When applied to an entire population with half in each group, Karin et al. [37] find, with some assumptions about the values of R_0 for each group, that with little mixing between the two groups, a 4-day work schedule, followed by a 10-day lockdown leads to a nearly steady decrease (small oscillations are present similar to Figure 16) in the number of infections in the population. Thus, people are working 4 days every other week. The model result is a consequence of the fact that anyone infected at work subsequently is isolated at home for most if not all of their infectious period, and of course would not return to work until recovered. The 4:10 cycle has an effective $R_0 < 1$ though it should be kept in mind that the model had little mixing between the two groups

D Daily testing at the start of the workday

In this section, we illustrate a five-fold reduction in exposure for daily tests administered at the start of the workday as opposed to at the end of the workday.

The crucial times are when the last negative test was taken, and when the first positive test is taken, and when it is reported. Take 0 to be the time when the last negative was taken, t_S (S for sick) when the first negative test result becomes known, and t_N for when the first negative test is taken. We assume the person becomes sick at some random time t_0 in the interval $[0, t_N]$.

Let $f_w(t)$ be 1 when the person is at work, and 0 when the person is not at work. Let $f_i(t - t_0)$ be infectivity of the person at time t . $f_i(t) = 0$ for $t \leq t_0$, and it increases monotonically to 1. (Eventually it goes back to 0, but we will only care about the time up to t_S when presumably the person is still infectious.)

Then the expected amount of time at work when the person is infectious is

$$\int_0^{t_N} 1/t_N \int_{t_0}^{t_S} f_i(t - t_0) f_w(t) dt dt_0. \quad (\text{D-21})$$

The inner integral is the exposure time if the person becomes infectious at t_0 , and the outer integral averages that over the time from 0 to t_N , by which time the person has become sick.

The first case has tests that are done in real time ($t_N = t_S = 1$) just before work and every day, measuring time in days. Further assume f_w is 1 for the first 8 hours (1/3 of a day), and 0 thereafter. Then the expected exposure time is

$$\int_0^{1/3} \int_{t_0}^{1/3} f_i(t - t_0) dt dt_0 = \int_0^{1/3} \int_0^{1/3-t_0} f_i(x) dx dt_0. \quad (\text{D-22})$$

If $f_i(x) = 1$ when $x > 0$, this integral is 1/18 of a day, which is 4/3 hours.

In the other scenario the tests are taken at the end of the work day, so $t_N = 1$, and reported at the beginning of the next work day, so $t_S = 5/3$,

once again assuming an 8 hour work day. Now f_w is 0 until the work day starts, which is overnight, when $t = 2/3$, and then it is 1 until the work day ends at $t = 1$, and then 0 again overnight. So the expected exposure time is

$$\int_0^1 \int_{t_0}^1 f_i(t - t_0) f_w(t) dt dt_0 = \int_0^1 \int_{\max(t_0, 2/3)}^1 f_i(t - t_0) dt dt_0. \quad (\text{D-23})$$

This breaks into two pieces,

$$\int_0^{2/3} \int_{2/3}^1 f_i(t - t_0) dt dt_0 + \int_{2/3}^1 \int_{t_0}^1 f_i(t - t_0) dt dt_0$$

If once again we assume $f_i(x) = 1$ when $x > 0$, the first integral is $2/9$, the second integral is $1/18$, and the total expected exposure time is $5/18$, which is 5 times as much as the first case.

If f_i doesn't jump to 1 instantaneously, the contrast is even worse. The first case uses the values of $f_i(x)$ for small x , up to $1/3$. The first integral in the second case uses the values of $f_i(x)$ for x bigger than $2/3$.

E Pooled testing

Suppose there is a probability I that a person is infected and so a probability $H = 1 - I \approx 1$ that a given person is healthy.

Suppose we need to test $N \gg 1$ people to determine who is infected. Further suppose the test is perfect and that a given person's health status is independent of that of the others.

Adopt the protocol that samples will be pooled into G groups of M members each so that $GM = N$. Further, if a given group tests positive, follow up tests will be done on each of the M member of the group.

The probability that a given group is positive (i.e., contains at least one infected individual) is $(1 - H^M)$, so that the total number of tests that need to be performed on average is

$$G[1 + M(1 - H^M)] \quad (\text{E-24})$$

where the first term are the pooled tests and the second are the M individual followup tests that need to be done if a group tests positive. Since $G = N/M$, the total number of tests that need to be done is

$$N[M^{-1} + (1 - H^M)]. \quad (\text{E-25})$$

Comparing this to the N individual tests that would need to be done if there were no pooling, we find the pooling efficiency ϵ to be given by

$$\epsilon^{-1} = [M^{-1} + (1 - H^M)] \quad (\text{E-26})$$

Determining the group size M^* that maximizes ϵ involves a transcendental equation. But for $IM \ll 1$, (i.e., it's rare to find a group with an infected member), there is a simplification: the factor in parentheses is

$$1 - H^M = 1 - (1 - I)^M \approx IM \quad (\text{E-27})$$

In that case,

$$\epsilon^{-1} = (M^{-1} + IM). \quad (\text{E-28})$$

The maximal efficiency ϵ^* is maximized at $M^* = I^{-\frac{1}{2}}$, so that $IM^* = I^{\frac{1}{2}}$ and $\epsilon^* = \frac{1}{2} I^{-\frac{1}{2}}$.

In the case that $I = 0.01$, the optimal group size is $M^* = 10$, the probability that a group tests positive is $IM^* = 0.1$, and the total number of tests required is reduced by a $\epsilon^* = 5$.

Some comments:

1. When I is larger than 0.01, pooled testing (at least for the protocol chosen) doesn't make much sense. If $I = 0.1$, then the optimal $M^* = 3$, and the efficiency is $\epsilon^* = 1.5$. At this level, the logistical challenges created by pooling do not justify the efficiency gain.
2. When I is very small, there appears to be a potential for a large benefit. If $I = 0.0001$, then $M^* = 100$ and $\epsilon^* = 50$. But, in fact, the groups can't get too large since a single positive would get too diluted to be detected reliably. Current PCR tests can handle up to $M = 10$, which is the optimum for $I = 0.01$.
3. Group testing is a mature subject https://en.wikipedia.org/wiki/Group_testing and no doubt there are protocols that are even more efficient than chosen here, but none simpler. They can be tailored to where the bottlenecks are (sample collection, insufficient reagent, preparation time, etc.).
4. It will make sense to choose groups that are in close contact (e.g., families living together or co-workers). If there's one infected, likely more are, so the signal would be amplified. And if the pooled sample tests positive, the whole group should be assumed to be exposed and so quarantined while individual testing is performed.

E.1 Incorporating test imperfections

A test will be characterized by P_D , the probability of detection (also called the sensitivity) and P_{FA} , the probability of a false alarm (a false positive). An imperfect test has $P_D < 1$ and finite, but hopefully small, P_{FA} .

For such an imperfect test, the average number of tests required for the protocol outlined above is

$$G[1 + M(1 - H^M)P_D + MH^M P_{FA}] \quad (\text{E-29})$$

Here, the first term is the initial test of the group, the second corresponds to having to test the whole group when there is a true positive, and the third is having to test the whole healthy group when there's a false alarm in the group test. Then the efficiency will be

$$\epsilon^{-1} = [M^{-1} + (1 - H^M)P_D + H^M P_{FA}] \quad (\text{E-30})$$

For $MI \ll 1$, this can be approximated by

$$\epsilon^{-1} = [M^{-1} + MI(P_D - P_{FA}) + P_{FA}] \quad (\text{E-31})$$

A crucial question is how P_D and P_{FA} depend upon dilution (i.e., the group size M). P_{FA} is the probability of a false alarm in the test of an entirely healthy group – it seems plausible that this will be independent of M . It also can't be too large, lest the test be essentially useless. On the other hand, P_D is the probability of detecting at least one infected individual in a sample diluted M -fold. Given the non-linear amplification of PCR, it's plausible to take P_D as independent of M up to some cutoff M_C , and then zero for larger values of M . Under these assumptions, the efficiency can be maximized as before to find

$$M^* = \min\{M_C, [I(P_D - P_{FA})]^{-\frac{1}{2}}\} \quad (\text{E-32})$$

and

$$\epsilon^* = [M^{*-1} + M^*I(P_D - P_{FA}) + P_{FA}]^{-1} \quad (\text{E-33})$$

It's reasonable to suppose that any test will be operated conservatively to put P_D close to 1 while tolerating some level of false alarms. As an example, Figures 17 and 18 below show how M^* and ϵ^* vary with I and P_{FA} when $M_C = 10$ and $P_D = 1$.

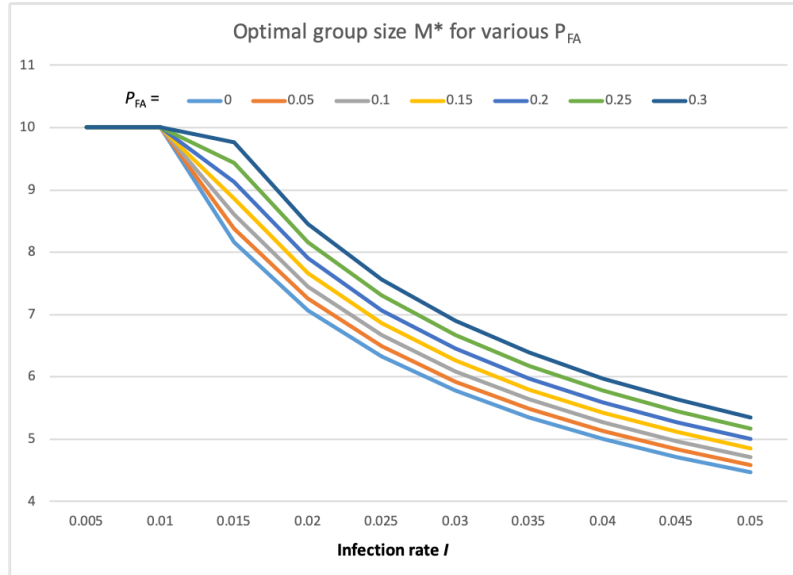


Figure 17: Optimal group size M^* for various P_{FA} ; $M_C = 10$ and $P_D = 1$.

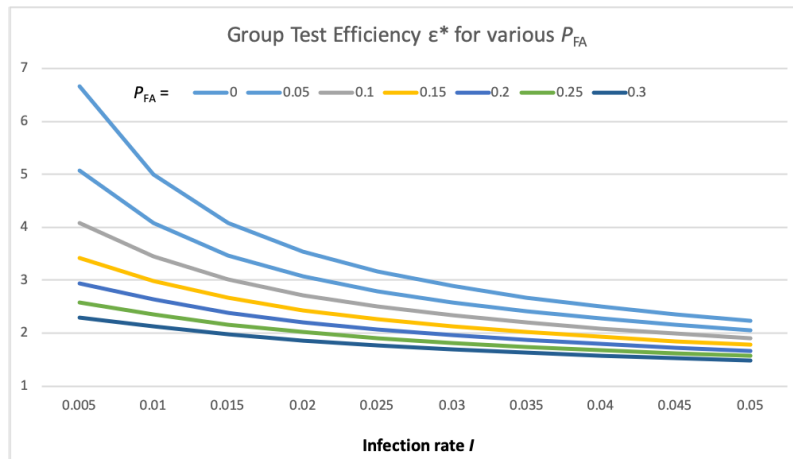


Figure 18: Optimal group size ϵ^* for various P_{FA} ; $M_C = 10$ and $P_D = 1$.

The optimal group size depends only weakly on the false alarm rate, while the group test efficiency degrades considerably as PFA increases. In sum, at least with the testing protocol as defined and for plausible false alarm

rates, group testing would be useful only for infection rates less than about 1%.

F Marginal risk from COVID-19

This brief note estimates the added risk of death due to COVID-19 as a function of age. The result depends on factors that are as yet not well established.

The Social Security administration publishes a tabulation of annual mortality risk vs. age for the U.S. population.⁸ The fraction of the population in each 1 year age bin that die per year, as a function of age, is plotted in Figure 19.

The Case Fatality Ratio (CFR) for COVID-19 remains a topic of ongoing investigation, complicated by the lack of full knowledge about the number of asymptomatic cases. The data are typically coarsely binned. Figure 20 shows the CFR estimates from CDC, along with a quadratic fit. The numbers used here are for symptomatic cases. If only half the infections produce symptoms then the marginal risk numbers below fall by a factor of two.

The rate at which the US population is contracting COVID-19 was drawn from data at Oxford University.⁹ They estimate that 1.6 million cases were contracted in the 60 day interval between April 1 and May 30, with a very linear trend. At this rate we should expect 9.7 million additional cases each year, which is a fraction of $9.7\text{E}6/320\text{E}6 = 3\%$ of the population infected per year.

Under the assumption that the rate of infection is age-independent, we can compare the mortality risk from COVID-19 (the infection rate times the CFR-fit in each age bracket) to that of all other causes.

Conclusions:

1. The marginal annual mortality risk increase at the current infection rate is a few percent, Fig 21.
2. Despite the well-publicized age-dependence of COVID-19 outcomes, this marginal risk is surprisingly age-independent.

⁸https://www.ssa.gov/oact/STATS/table4c6_2016.html#fn1

⁹<https://ourworldindata.org/mortality-risk-covid>

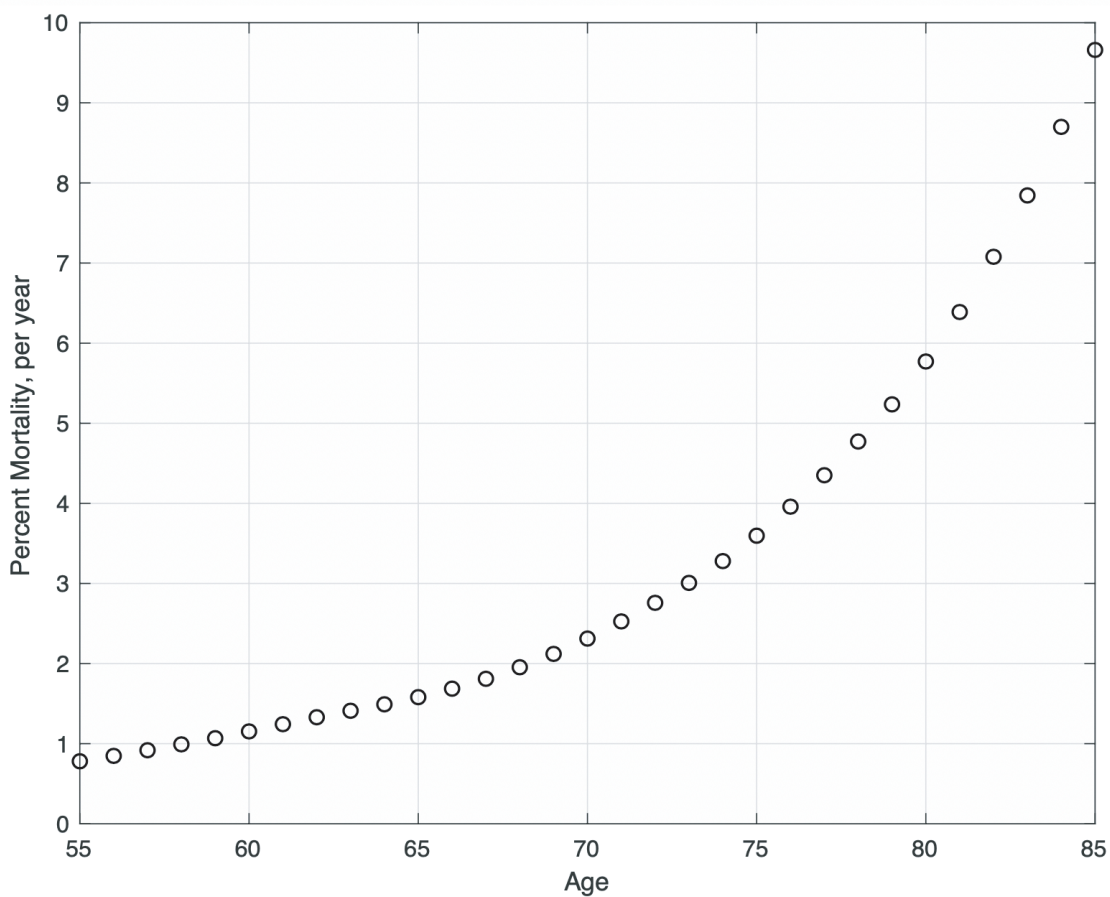


Figure 19: Mortality (fraction of cohort that dies, per year) as a function of age. From Social Security Administration.

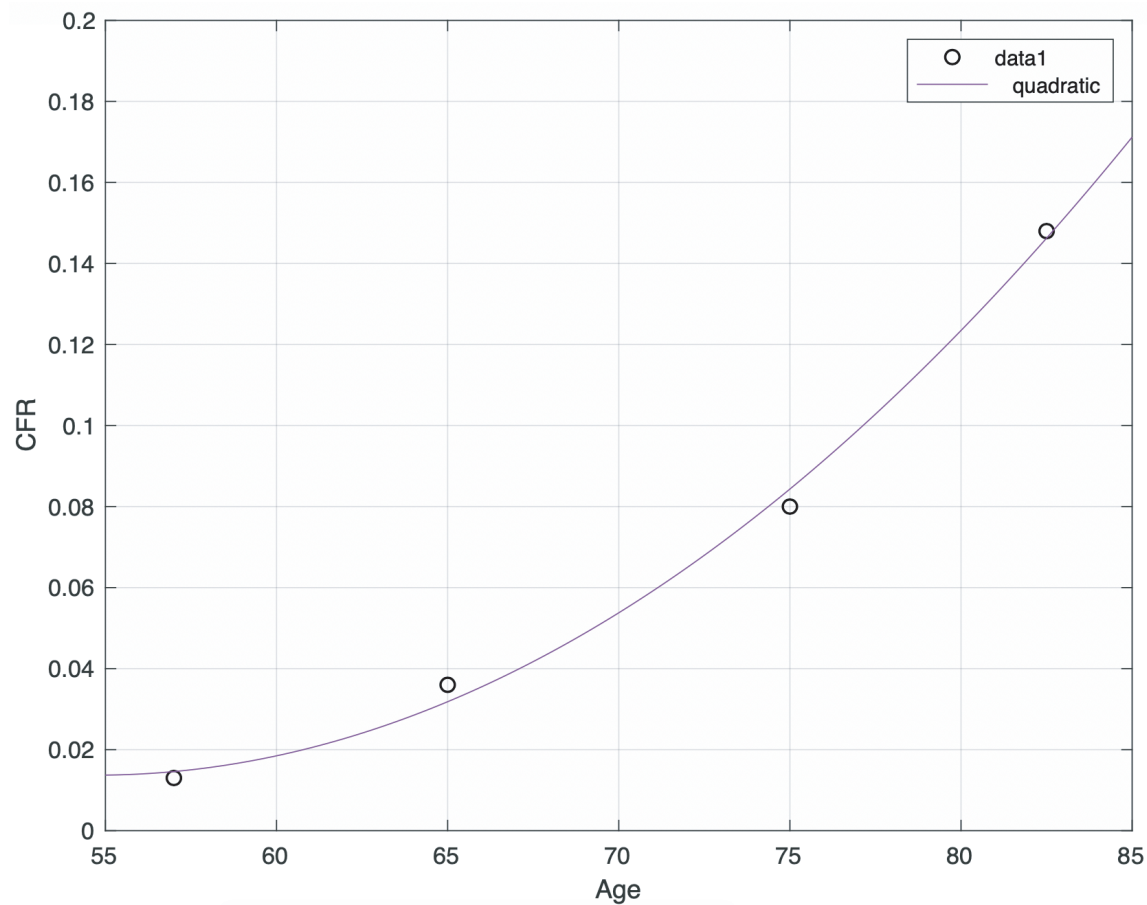


Figure 20: Case Fatality Ratio (CFR) estimates from CDC, for symptomatic cases, along with a quadratic fit.

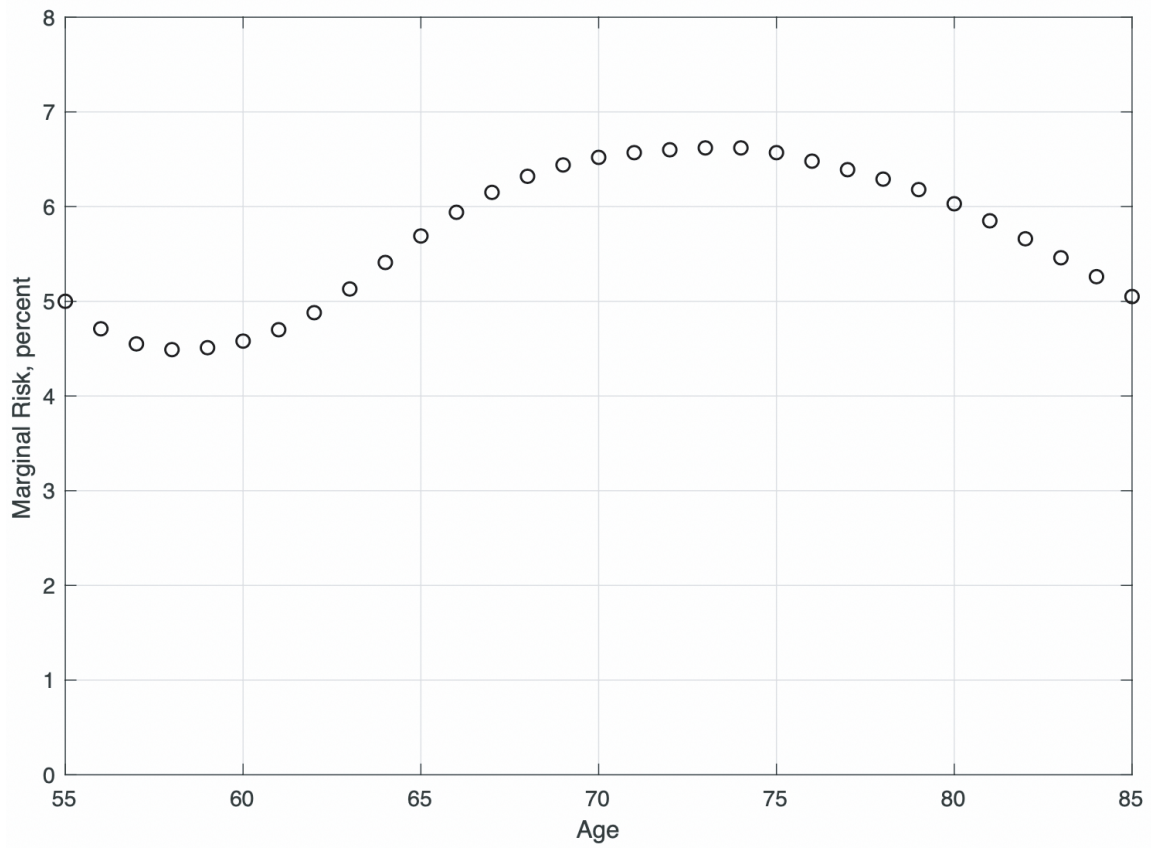


Figure 21: Marginal risk from COVID as a function of age. If we include asymptomatic cases in the normalization of the CFR, the marginal risk is about a factor of two smaller.

G Expected utility for decision-making

Aerosol transmission may be an important part of the way the SARs-CoV-2 virus spreads. Evans [19] and Section 3 provide a means of assessing relative risk to aerosol transmission for those working or learning in labs with fresh air changes or filtering. This note provides a method for computing the relative risk for different numbers of people in a room and an example using the method of Evans [19] in order to assess risk as a function of parameter variations. A research group has N people and the PI must decide how many, u , will work at the same time in a laboratory. f is the poorly known infection rate among the N researchers in the group and t_o is the time between infection and the appearance of symptoms. Reducing the number of transmissions of the SARS-CoV-2 virus between group members while working in the lab is the PI's goal. The time rate of transmission in the lab, g may be calculated using the Evans [19] method and depends on air-changes in the room, the virus lifetime as an aerosol, and so on. A workday in the lab is τ .

If $u = 0$ or $u = 1$, transmission cannot take place. If the PI chooses to have two people in the room, then the probability none are infected is $(1 - ft_o)^2$, that one is infection $2(1 - ft_o)ft_o$, and that both are infected $(ft_o)^2$. Transmission occurs if one researcher is infected and the other is not and the probability is $g\tau$. The combined probability for transmission to occur if the PI chooses to have two people in the lab is $2(1 - ft_o)ft_og\tau$ for one day of work.

In general,

$$P(u \text{ chosen} | m \text{ infected}) = P(u | m) = \binom{u}{m} ft_o^m (1 - ft_o)^{u-m}.$$

With m infected and $u - m$ uninfected, there are $m(u - m)$ ways transmission can occur and $g\tau(1 - g\tau)^{n-1}m(u - m)$ gives the probability for one transmission. If $ug\tau \ll 1$, then two transmissions are unlikely compared to one transmission and the probability is $g\tau m(u - m)$:

$$\begin{aligned} P(\geq 1 \text{ transmission } u \text{ chosen} | m \text{ infected}) &= P(1, u | m) \\ &= \binom{u}{m} ft_o^m (1 - ft_o)^{u-m} g\tau m(u - m) \end{aligned}$$

No. in room, u	$P_{rel}(1, u)$	$P(1, u)$	
		550 seat lecture	400 sq. ft. lab
2	1	2.6×10^{-6}	0.00013
3	3	7.7×10^{-6}	0.00038
4	6	0.00015	0.00076
5	10	0.000025	0.00113
10	45	0.00011	0.0057
30	435	0.0011	0.055
100	4,950	0.012	> 1 transmission

Table 2: Relative and absolute infection rates per hour. For the absolute probabilities, the number gives the probability of one transmission per hour, > 1 transmission means more than one transmission per hour is likely.

and summing over m gives for one or more transmissions between u people,

$$P(1, u) = \sum_{m=1}^{u-1} \binom{u}{m} f t_o^m (1 - f t_o)^{u-m} g \tau m (u - m) \quad (\text{G-34})$$

Though important parameters, f , t_o , and g remain uncertain, Eq. G-34 still gives an important result for the probability of transmission relative to the probability of transmission with just two people in the room,

$$P_{rel}(1, u) = \frac{P(1, u)}{P(1, 2)} = \frac{1}{2} u (u - 1) \quad (\text{G-35})$$

Table 2 gives the increased relative risk as people are added to the room.

The absolute probability $P(1, 2)$ depends on the infection rate between two people g , the infection rate of the research group f , and t_o , the time between infection and symptoms. In the state of Massachusetts, on May 12, there were 1,000 new cases per day for a population of 6.7 million. The daily new cases result from testing, giving a lower limit of $f > 0.00015/\text{day}$, as many infections go untested and unreported. $t_o = 5.1 \pm 0.7$ days [17]. Evans [19] Eqs. 1 and 8 gives g for a specific room. Table 3 gives the model parameters for a typical lab and 500 seat lecture room.

Table 2 gives the absolute probabilities per hour of exposure for a lab and lecture hall in the right columns. If researchers in a lab do not have much contact outside of the group for ≈ 2 weeks or 80 hours, confidence builds that they are not and unlikely to become infected. For five people

Parameter	Lab Room 400 sq. ft.	Lecture Room 551 seats, 5,728 sq. ft.
r_{src}	1 nL/min	1 nL/min
r_{room}	445 m ³ /h	21,000 m ³ /h
t	1 h	1 h
g	0.078/h	0.0025/h
$P(1, 2)$	0.00013	2.6×10^{-6}

Table 3: Model parameters for a typical lab room and large lecture room for Evans [19] Eq. 8.

working together for 80 hours have a group probability of 9% of becoming infected. For a large lecture hall, a term of lectures is 42 hours of lecture (14 weeks, 3 hours per week), so a group of 100 people have a probability of 40% of transmitting one infection between two members.

Knowing the probability of transmitting a single COVID-19 infection, to create one additional case, does not help decision makers very much. They need to balance the damage from an additional case against the gain from operating in an environment that allows a case to be transmitted. This memo lays out one method of balancing the gain and loss.

Ernst Weber (1795-1878) and his student Gustav Fechner (1801-1887) were pioneers of psychophysics – the study of human perception of physical stimulus. Their body of work has been applied to the wider realm of human endeavor including perception, finance, and numerical cognition. The Weber-Fechner law relates perception p with stimulus S ,

$$p = k \ln \frac{S}{S_o}$$

where S_o is a reference stimulus value. For $S = S_o + \delta$ and $\delta \ll S_o$, $p \sim k\delta/S_o$. k is a constant.

Daniel Bernoulli (1700-1782) employed similar ideas to the St. Petersburg Paradox,

“The determination of the value of an item must not be based on the price, but rather on the utility it yields. There is no doubt that a gain of one thousand ducats is more significant to the pauper than to a rich man though both gain the same amount.”

to lay the groundwork for expected utility theory. Expected utility theory takes into account the human perception of gain or loss as proportional rather than absolute. This work uses the same notion.

The following is a calculation of expected utility of *adding* another lab of the type described in above. A number u of researchers share the lab for 80 hours, two work weeks, after which time we could conclude that none are infected and will not become infected if they follow proper safety procedures. If one of the researchers turns out to be infected, the probability of one COVID-19 transmission between two researchers during their 80 hours together is,

$$p_{trans} = \frac{1}{2}u(u-1)p(1,2)$$

and $p(1,2) = 0.00013$ for a set reference parameters in [19]. For u researchers, the expected gain of being able to work together in the lab is $\Delta g = 80u$ hours. Over the same two weeks, the total research capacity of a medium sized research university (10,000 faculty, staff, and student researchers) operating at 25% capacity is $g = 196,400$ hours. If, as a result of the u researchers working together, an infection occurs, there are two losses: $\Delta l_1 = 80$ hours of lost research, out of $l_1 = 196,400$ hours and the personal time lost to the researcher. If the researcher is 30 years old, they have a life expectancy of $l_2 = 47$ years and will lose $\Delta l_2 = 336$ hours of their life, assuming a two week course of COVID-19. The expected utility is then,

$$E(u) = (1-p)(u) \ln \frac{g + \Delta g}{g} - p \ln \left(\frac{l_1 + \Delta l_1}{l_1} + \frac{l_2 + \Delta l_2}{l_2} \right).$$

Fig. 22 shows $E(u)$ as a function of u for different fractions of $p(1,2) = 0.00013$: $p(1,2)$ has to be reduced by a factor of 3 for two people working in the lab to yield a positive return and a reduction of 30 yields would have 10 people working in the lab as a maximum expected utility.

The analysis above also considers a 550 seat lecture room with a $p(1,2) = 2.6 \times 10^{-6}$. The gain in this case comes from the exposure of u students to a term of lectures (14 weeks, 3h per week), $\Delta g = 42u$. A medium sized university with 4,350 students each taking one large lecture course will have $g = 190,260$ student lecture hours.

Three loss terms occur in the case of an infection: the loss of a term of instruction in the case of an infection, $\Delta l_1 = 42$, $l_1 = 761,040$, the loss of $\Delta l_2 = 336$ hours of health by a 20 year old student with a life expectancy of $l_2 = 490,560$ hours=56 years and, with probability of $1/u$, the 60% marginal probability of the loss of the life of a 60 year lecturer, whose non-COVID-19

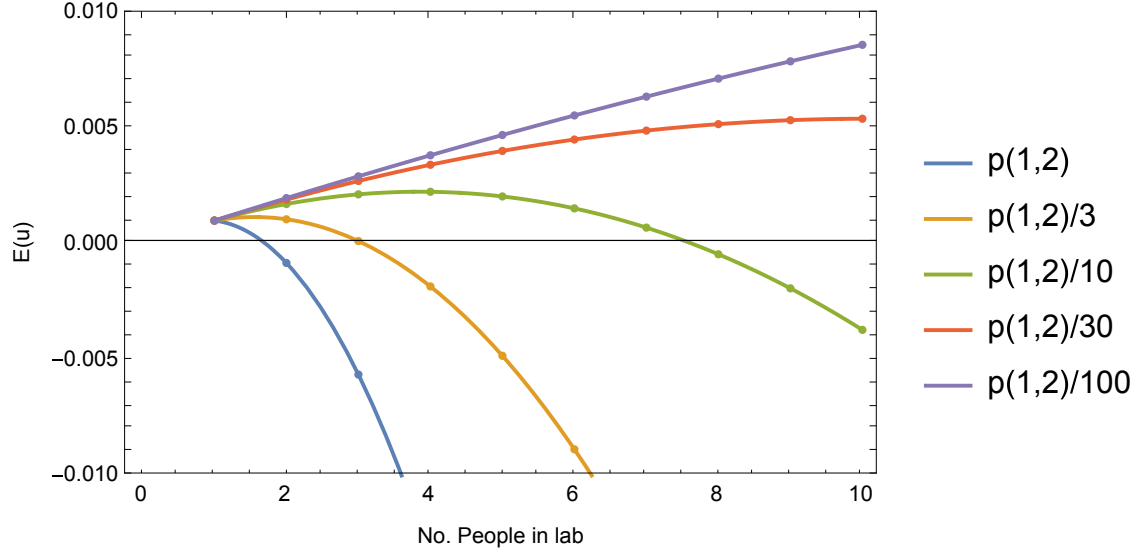


Figure 22: Expected utility as a function of the number of people in the lab for an 80 hours. The dots indicate integral numbers of people.

annual probability of death is $l_3 = 0.0115^{10}$. Fig. 23 and Fig. 24 shows that at nominal values, the expected utility is largest with 8 people in the lecture room, and if $p(1,2)$ can be reduced by a factor of 10, about 66 people can be in the lecture room. Fig. 24 shows the threat to the life of the lecturer does not dominate the expected utility owing to the $1/u$ probability that the lecturer is the one on the receiving end of the infection transmission.

These two examples considered the incremental increase in the expected utility from adding a lab or lecture and the increment is small because $\Delta g \ll g$. A lab or campus wide optimization over all venues E_i each populated by u_i people and maximizing,

$$E(u_1, \dots, u_n) = \sum_{i=1}^n E_i(u_1, \dots, u_n)$$

to find the optimal set of venue populations u_i . E_i embodies the physical characteristics of the i^{th} lab or lecture room and depends $u_1, \dots, u_n$ through the total gain or loss. The process could be tiered by defining set of top priority rooms and optimizing for them first, then optimizing for a second tier and so on.

¹⁰Appendix F gives a non-COVID-19 mortality rate of 0.011 and a case fatality rate of 0.018 for a 60 year old, so given that a 60 year contracts COVID-19, their mortality rate increases by 60%.

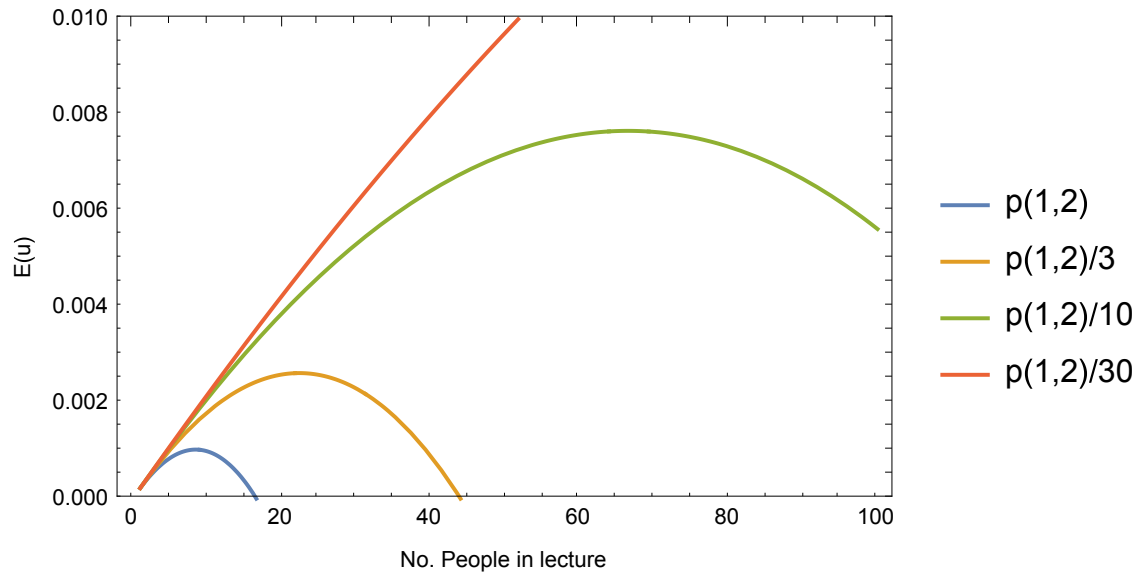


Figure 23: Expected utility as a function of the number of people in a 551 seat lecture hall for different values of $p(1,2)$.

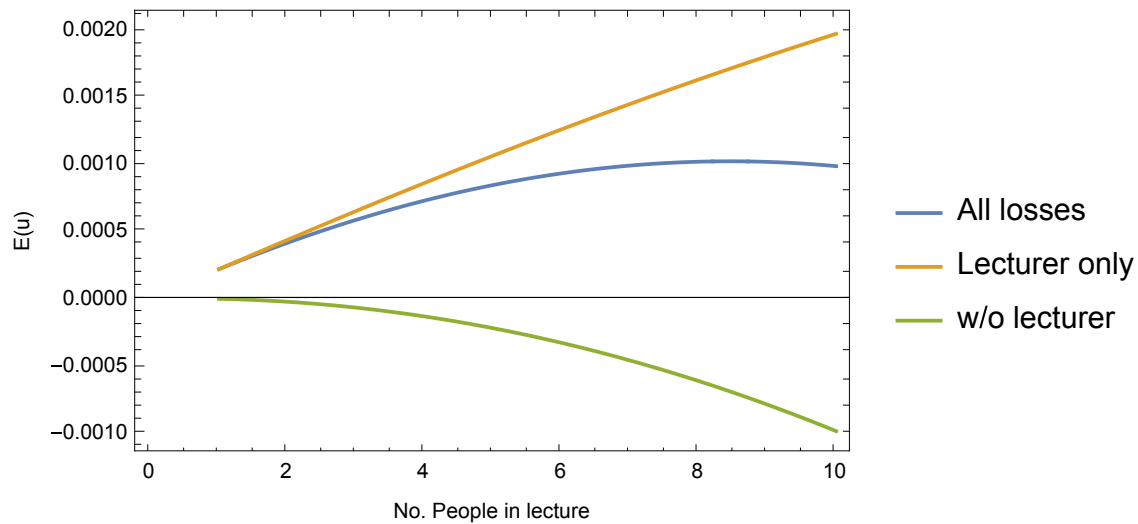


Figure 24: Expected utility as a function of the number of people in a 551 seat lecture hall, breaking out losses from the lecturer only (orange), everyone aside from the lecturer (green) and total (blue).

The loss and gain may be separately weighted,

$$E(u) = \alpha(1-p)(u) \ln \frac{g + \Delta g}{g} - \beta p \ln \left(\frac{l_1 + \Delta l_1}{l_1} + \frac{l_2 + \Delta l_2}{l_2} \right).$$

to reflect that the gain in productivity may be valued less than the loss from a transmission of the disease, $\alpha < \beta$. Determining the weights before carrying out the calculation reduces the bias in the relative valuation of the gain and loss.

H Simple analysis of the impact of false-positive tests

This section considers the using of testing at low disease prevalence ($<10\%$) when the test has a high false-positive rate.

N is the population size, f is the fraction of the population who have been infected but not yet showing symptoms, g is the fraction of the population that can be tested at any one time, and h is the fraction of tests that return positive for an uninfected patient. The false positive tests follows a Poisson distribution.

The number of infected people is Nf and Nfg is the number of true positives. The number of false positive tests is Nh . Then $s = Nfg/\sqrt{Nh} = fg\sqrt{N/h}$ is the significance of the false positive tests – the number of standard deviations above the mean background of Nh false positives tests that Nfg tests lie. For a specified significance, $h = Nf^2g^2/s^2$ gives the false positive rate that may be tolerated while still yielding a specified significance s . Figure 25 provides an example for $N = 10,000$.

Figure 25 illustrates how false positives affect the value of a test. First we consider the case with a university with 10,000 researchers and 1% of its population infected but pre-symptomatic. If the university is able to test 10% of its population (green square), the contour indicates that a test false-positive rate of no more than 0.3% can be tolerated before the true-positive rate becomes less than one standard deviation above the expected false positive rate.

A second scenario considers the university is able to test 55% of its population by using a different test having a false-positive rate of 10%. The infection rate is still 1% (red star on figure), but the figure shows the testing program will yield a significant infection signal as a 30% false positive rate can be tolerated.

Finally, things take a turn for the worst and the university has 2.5% infection rate (blue circle on figure), and the university achieves a 55% testing rate. In this case the true positives swamp the false positive results.

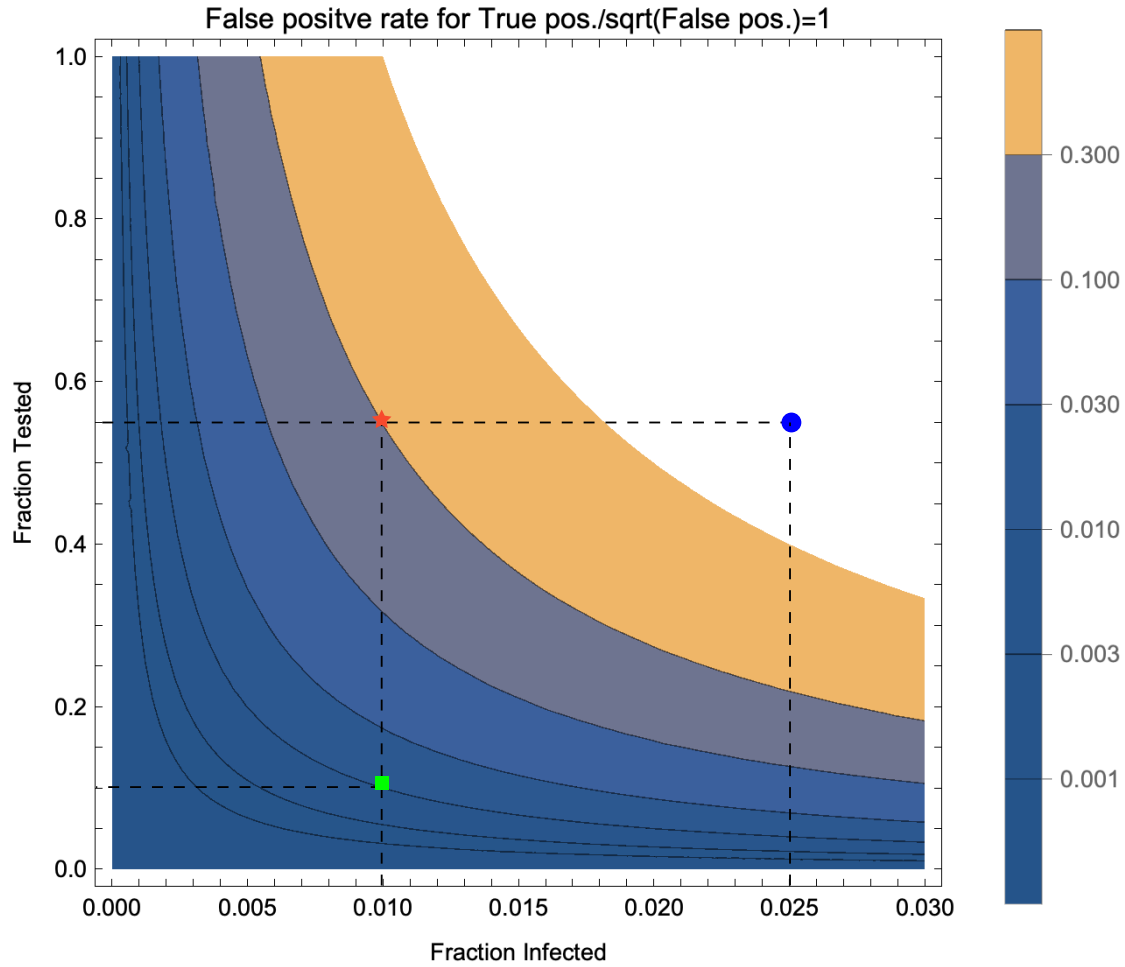


Figure 25: Contour plot for $N = 10,000$ tests. The contours are lines of constant false positive fraction h . The star is an example showing that for a $f=1\%$ infection rate and a test fraction of g of 56%, a false positive rate of $h=30\%$ less is required for a true positive test signal to lie one standard deviation above background ($s = 1$). Other symbols for scenarios explained in the main text.

References

- [1] Benjamin M. Althouse, Edward A. Wenger, Joel C. Miller, Samuel V. Scarpino, Antoine Allard, Laurent Hébert-Dufresne, and Hao Hu. Stochasticity and heterogeneity in the transmission dynamics of SARS-CoV-2, 2020.
- [2] Sima Asadi, Anthony S. Wexler, Christopher D. Cappa, Santiago Barreda, Nicole M. Bouvier, and William D. Ristenpart. Aerosol emission and superemission during human speech increase with voice loudness. *Scientific Reports*, 9(1):2348, December 2019. ISSN 2045-2322. doi: 10.1038/s41598-019-38808-z.
- [3] Augenbraun, Lasner, Raval Sawaoka Mitra, Prabhu, and J Doyle. Assessment and Mitigation of Aerosol Airborne SARS-CoV-2 Transmission in Laboratory and Office Environments. *Journal of Occupational Environmental Hygiene*, March 2020.
- [4] Linlin Bao, Wei Deng, Hong Gao, Chong Xiao, Jiayi Liu, Jing Xue, Qi Lv, Jiangning Liu, Pin Yu, Yanfeng Xu, Feifei Qi, Yajin Qu, Fengdi Li, Zhiguang Xiang, Haisheng Yu, Shuran Gong, Mingya Liu, Guanpeng Wang, Shunyi Wang, Zhiqi Song, Wenjie Zhao, Yunlin Han, Linna Zhao, Xing Liu, Qiang Wei, and Chuan Qin. Reinfection could not occur in SARS-CoV-2 infected rhesus macaques. Preprint, Microbiology, March 2020.
- [5] Yinon M. Bar-On, Ron Sender, Avi I. Flamholz, Rob Phillips, and Ron Milo. A quantitative compendium of COVID-19 epidemiology. *arXiv:2006.01283 [q-bio]*, June 2020.
- [6] Bernhard Blythe. St. louis saw the deadly 1918 spanish flu epidemic coming. shutting down the city saved countless lives. *St. Louis Post Dispatch*, March 2020.
- [7] Paul Jacob Bueno de Mesquita, Catherine J. Noakes, and Donald K. Milton. Quantitative aerobiologic analysis of an influenza human challenge-transmission trial. *Indoor Air*, page ina.12701, June 2020. ISSN 0905-6947, 1600-0668. doi: 10.1111/ina.12701.
- [8] Emma Cave. COVID-19 Super-spreaders: Definitional Quandaries and Implications. *Asian Bioethics Review*, 2020. doi: 10.1007/s41649-020-0018.
- [9] C. Y.H. Chao, M. P. Wan, and G. N. Sze To. Transport and Removal of Expiratory Droplets in Hospital Ward Environment. *Aerosol Science*

- and Technology*, 42(5):377–394, March 2008. ISSN 0278-6826, 1521-7388. doi: 10.1080/02786820802104973.
- [10] C.Y.H. Chao, M.P. Wan, L. Morawska, G.R. Johnson, Z.D. Ristovski, M. Hargreaves, K. Mengersen, S. Corbett, Y. Li, X. Xie, and D. Katoshevski. Characterization of expiration air jets and droplet size distributions immediately at the mouth opening. *Journal of Aerosol Science*, 40(2):122–133, February 2009. ISSN 00218502. doi: 10.1016/j.jaerosci.2008.10.003.
- [11] Jun Chen, Tangkai Qi, Li Liu, Yun Ling, Zhiping Qian, Tao Li, Feng Li, Qingnian Xu, Yuyi Zhang, Shuibao Xu, et al. Clinical progression of patients with COVID-19 in Shanghai, China. *Journal of Infection*, 2020.
- [12] A.N Cohen and B. Kessel. False positives in reverse transcription PCR testing for SARS-CoV-2. *medRxiv*, 2020. doi: 10.1101/2020.04.26.20080911. URL <https://www.medrxiv.org/content/early/2020/05/20/2020.04.26.20080911>.
- [13] Talib Dbouk and Dimitris Drikakis. On respiratory droplets and face masks. *Physics of Fluids*, 32(6):063303, June 2020. ISSN 1070-6631, 1089-7666. doi: 10.1063/5.0015044.
- [14] Y. Drossinos and N.I. Stilianakis. What aerosol physics tells us about airborne pathogen transmission. *Aerosol Science and Technology*, 54(6): 639–643, 2020.
- [15] D. A. Edwards, J. C. Man, P. Brand, J. P. Katstra, K. Sommerer, H. A. Stone, E. Nardell, and G. Scheuch. Inhaling to mitigate exhaled bioaerosols. *Proceedings of the National Academy of Sciences*, 101(50): 17383–17388, December 2004. ISSN 0027-8424, 1091-6490. doi: 10.1073/pnas.0408159101.
- [16] S.E. Eikenberry et al. To mask or not to mask: Modeling the potential for face mask use by the general public to curtail the covid-19 pandemic, 2020. URL <https://doi.org/10.1101/2020.04.06.20055624>.
- [17] S.A. Lauer et al. The incubation period of coronavirus disease 2019 (covid-19) from publicly reported confirmed cases: estimation and application. *preprint*, 2020.
- [18] Xiaobo Yang et al. Clinical course and outcomes of critically ill patients with SARS-CoV-2 pneumonia in Wuhan, China: a single-centered, retrospective, observational study. *The Lancet Respiratory Medicine*, 8(5), 2020. doi: 10.1016/S2213-2600(20)30079-5.

- [19] Matthew Evans. Avoiding COVID-19: Aerosol Guidelines. Preprint, Infectious Diseases (except HIV/AIDS), May 2020.
- [20] S. Flaxman, S. Mishra, A. Gandy, et al. Estimating the effects of non-pharmaceutical interventions on COVID-19 in Europe. *Nature*, 2020. doi: 10.1038/s41586-020-2405-7.
- [21] Huizhi G. Wong j. Xiao J. Shiu E. Ryu S. Fong, Min and B. Cowling. Nonpharmaceutical measures for pandemic influenza in nonhealth-care settings— social distancing measures, emerging infectious diseases. Technical report.
- [22] Center for Disease Control and Prevention. Interim Clinical Guidance for Management of Patients with Confirmed Coronavirus Disease (COVID-19); updated 2 June 2020. Technical report, 2020. URL <https://www.cdc.gov/coronavirus/2019-ncov/hcp/clinical-guidance-management-patients.html>.
- [23] National Centre for Infectious Diseases and Academy of Medicine the Chapter of Infectious Disease Physicians. Position Statement from the National Centre for Infectious Diseases and the Chapter of Infectious Disease Physicians, Academy of Medicine, Singapore – 23 May 2020. Technical report, 2020. URL [https://www.ams.edu.sg/view-pdf.aspx?file=media%5c5556_fi_331.pdf&ofile=Period+of+Infectivity+Position+Statement+\(final\)+23-5-20+\(logos\).pdf](https://www.ams.edu.sg/view-pdf.aspx?file=media%5c5556_fi_331.pdf&ofile=Period+of+Infectivity+Position+Statement+(final)+23-5-20+(logos).pdf).
- [24] Monica Gandhi, Deborah S Yokoe, and Diane V Havlir. Asymptomatic transmission, the Achilles’ heel of current strategies to control COVID-19, 2020.
- [25] Atul Gawande. Amid the Coronavirus Crisis, a Regimen for Reentry. *The New Yorker*, May 2020.
- [26] Katelyn Gostic, Ana CR Gomez, Riley O Mummah, Adam J Kucharski, and James O Lloyd-Smith. Estimated effectiveness of symptom and risk screening to prevent the spread of COVID-19. *Elife*, 9:e55570, 2020.
- [27] Mecher C. E. Hatchett, Richard and Lipsitch M. Public health interventions and epidemic intensity during the 1918 influenza pandemic. *Pubs. NAS*, 104:7582–7587, 2007.
- [28] Xi He, Eric H. Y. Lau, Peng Wu, Xilong Deng, Jian Wang, Xinxin Hao, Yiu Chung Lau, Jessica Y. Wong, Yujuan Guan, Xinghua Tan, Xiaoneng Mo, Yanqing Chen, Baolin Liao, Weilie Chen, Fengyu Hu, Qing Zhang, Mingqiu Zhong, Yanrong Wu, Lingzhai Zhao, Fuchun Zhang,

- Benjamin J. Cowling, Fang Li, and Gabriel M. Leung. Temporal dynamics in viral shedding and transmissibility of COVID-19. *Nature Medicine*, 26(5):672–675, 2020. doi: 10.1038/s41591-020-0869-5. URL <https://doi.org/10.1038/s41591-020-0869-5>.
- [29] Xi He, Eric H. Y. Lau, Peng Wu, Xilong Deng, Jian Wang, Xinxin Hao, Yiu Chung Lau, Jessica Y. Wong, Yujuan Guan, Xinghua Tan, Xiaoneng Mo, Yanqing Chen, Baolin Liao, Weilie Chen, Fengyu Hu, Qing Zhang, Mingqiu Zhong, Yanrong Wu, Lingzhai Zhao, Fuchun Zhang, Benjamin J. Cowling, Fang Li, and Gabriel M. Leung. Temporal dynamics in viral shedding and transmissibility of COVID-19. *Nature Medicine*, 26(5):672–675, 2020. doi: 10.1038/s41591-020-0869-5. URL <https://doi.org/10.1038/s41591-020-0869-5>.
- [30] S. Hsiang et al. The effect of large-scale anti-contagion policies on the COVID-19 pandemic. *Nature*, 2020. doi: 10.1038/s41586-020-2404-8.
- [31] David S. Hui, Benny K. Chow, Leo Chu, Susanna S. Ng, Nelson Lee, Tony Gin, and Matthew T. V. Chan. Exhaled Air Dispersion during Coughing with and without Wearing a Surgical or N95 Mask. *PLoS ONE*, 7(12):e50845, December 2012. ISSN 1932-6203. doi: 10.1371/journal.pone.0050845.
- [32] Roger W. Humphry, Angus Cameron, and George J. Gunn. A practical approach to calculate sample size for herd prevalence surveys. *Preventive Veterinary Medicine*, 65(3-4):173–188, October 2004. ISSN 0167-5877. doi: 10.1016/j.prevetmed.2004.07.003.
- [33] Guanmin Jiang, Xiaoshuai Ren, Yan Liu, Hongtao Chen, Wei Liu, Zhaowang Guo, Yaqin Zhang, Chaoqun Chen, Jianhui Zhou, Qiang Xiao, and Hong Shan. Application and optimization of RT-PCR in diagnosis of SARS-CoV-2 infection. Preprint, Infectious Diseases (except HIV/AIDS), February 2020.
- [34] G.R. Johnson, L. Morawska, Z.D. Ristovski, M. Hargreaves, K. Mengersen, C.Y.H. Chao, M.P. Wan, Y. Li, X. Xie, D. Katoshevski, and S. Corbett. Modality of human expired aerosol size distributions. *Journal of Aerosol Science*, 42(12):839–851, December 2011. ISSN 00218502. doi: 10.1016/j.jaerosci.2011.07.009.
- [35] Graham Richard Johnson and Lidia Morawska. The Mechanism of Breath Aerosol Formation. *Journal of Aerosol Medicine and Pulmonary Drug Delivery*, 22(3):229–237, September 2009. ISSN 1941-2711, 1941-2703. doi: 10.1089/jamp.2008.0720.

- [36] Terry C Jones, Barbara Mühlemann, Talitha Veith, Guido Biele, Marta Zuchowski, Jörg Hoffmann, Angela Stein, Anke Edelmann, Victor Max Corman, and Christian Drosten. An analysis of SARS-CoV-2 viral load by patient age. *medRxiv*, 2020. doi: 10.1101/2020.06.08.20125484. URL <https://www.medrxiv.org/content/early/2020/06/09/2020.06.08.20125484>.
- [37] Omer Karin, Yinon M. Bar-On, Tomer Milo, Itay Katzir, Avi Mayo, Yael Korem, Boaz Dudovich, Eran Yashiv, Amos J. Zehavi, Nadav Davidovich, Ron Milo, and Uri Alon. Adaptive cyclic exit strategies from lockdown to suppress COVID-19 and allow economic activity. Preprint, *Epidemiology*, April 2020.
- [38] Seong Eun Kim, Hae Seong Jeong, Yohan Yu, Sung Un Shin, Soosung Kim, Tae Hoon Oh, Uh Jin Kim, Seung-Ji Kang, Hee-Chang Jang, Sook-In Jung, and Kyung-Hwa Park. Viral kinetics of SARS-CoV-2 in asymptomatic carriers and presymptomatic patients. *International Journal of Infectious Diseases*, 95:441 – 443, 2020. ISSN 1201-9712. doi: <https://doi.org/10.1016/j.ijid.2020.04.083>. URL <http://www.sciencedirect.com/science/article/pii/S120197122030299X>.
- [39] Lauren M. Kucirka, Stephen A. Lauer, Oliver Laeyendecker, Denali Boon, and Justin 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*, 2020. doi: 10.7326/M20-1495. URL <https://doi.org/10.7326/M20-1495>. PMID: 32422057.
- [40] Stephen A. Lauer, Kyra H. Grantz, Qifang Bi, Forrest K. Jones, Qulu Zheng, Hannah R. Meredith, Andrew S. Azman, Nicholas G. Reich, and Justin Lessler. The Incubation Period of Coronavirus Disease 2019 (COVID-19) From Publicly Reported Confirmed Cases: Estimation and Application. *Annals of Internal Medicine*, 172(9):577–582, 2020. doi: 10.7326/M20-0504. URL <https://doi.org/10.7326/M20-0504>. PMID: 32150748.
- [41] Xu Li and Xiaochun Ma. Acute respiratory failure in COVID-19: is it “typical” ARDS? *Critical Care*, 24(198), 2020. doi: 10.1186/s13054-020-02911-9.
- [42] Yuguo Li, Hua Qian, Jian Hang, Xuguang Chen, Ling Hong, Peng Liang, Jiansen Li, Shenglan Xiao, Jianjian Wei, Li Liu, and Min Kang. Evidence for probable aerosol transmission of SARS-CoV-2 in a poorly ventilated restaurant. Preprint, *Infectious Diseases (except HIV/AIDS)*, April 2020.

- [43] Q. Long, B. Liu, and H. et al. Deng. Antibody responses to SARS-CoV-2 in patients with COVID-19. *Nature Medicine*, 2020. doi: 10.1038/s41591-020-0897-1.
- [44] Lei Luo, Dan Liu, Xin-long Liao, Xian-bo Wu, Qin-long Jing, Jia-zhen Zheng, Fang-hua Liu, Shi-gui Yang, Bi Bi, Zhi-hao Li, Jian-ping Liu, Wei-qi Song, Wei Zhu, Zheng-he Wang, Xi-ru Zhang, Pei-liang Chen, Hua-min Liu, Xin Cheng, Miao-chun Cai, Qing-mei Huang, Pei Yang, Xin-fen Yang, Zhi-gang Huang, Jin-ling Tang, Yu Ma, and Chen Mao. Modes of contact and risk of transmission in COVID-19 among close contacts. *medRxiv*, 2020. doi: 10.1101/2020.03.24.20042606. URL <https://www.medrxiv.org/content/early/2020/03/26/2020.03.24.20042606>.
- [45] Steven R. Lustig, John J. H. Biswakarma, Devyesh Rana, Susan H. Tilford, Weike Hu, Ming Su, and Michael S. Rosenblatt. Effectiveness of Common Fabrics to Block Aqueous Aerosols of Virus-like Nanoparticles. *ACS Nano*, 14(6):7651–7658, June 2020. ISSN 1936-0851, 1936-086X. doi: 10.1021/acsnano.0c03972.
- [46] Harvey B. Lipman PhD J. Alexander Navarro Alexandra Sloan Joseph R. Michalsen Alexandra Minna Stern Markel, Howard and Martin S. Cetron. Nonpharmaceutical interventions implemented by us cities during the 1918-1919 influenza pandemic. *J. Amer. Med. Asso.*, 298.
- [47] C. Menni, A.M. Valdes, and M.B. Freidin et al. Real-time tracking of self-reported symptoms to predict potential COVID-19. *Nature Medicine*, 2020. doi: 10.1038/s41591-020-0916-2.
- [48] Donald K. Milton, M. Patricia Fabian, Benjamin J. Cowling, Michael L. Grantham, and James J. McDevitt. Influenza Virus Aerosols in Human Exhaled Breath: Particle Size, Culturability, and Effect of Surgical Masks. *PLoS Pathogens*, 9(3):e1003205, March 2013. ISSN 1553-7374. doi: 10.1371/journal.ppat.1003205.
- [49] L. Morawska, G.R. Johnson, Z.D. Ristovski, M. Hargreaves, K. Mengersen, S. Corbett, C.Y.H. Chao, Y. Li, and D. Katoshevski. Size distribution and sites of origin of droplets expelled from the human respiratory tract during expiratory activities. *Journal of Aerosol Science*, 40(3):256–269, March 2009. ISSN 00218502. doi: 10.1016/j.jaerosci.2008.11.002.
- [50] S.S. Morse, R.L. Garwin, and P.J. Olsiewski. Next flu pandemic: What to do until the vaccine arrives. *Science*, 314:929, 2016.

- [51] A. Mueller et al. Assessment of fabric masks as alternatives to standard surgical masks in terms of particle filtration efficiency, April 2020. <https://doi.org/10.1101/2020.04.17.20069567>.
- [52] Arash Naseri, Omid Abouali, and Goodarz Ahmadi. Effect of turbulent thermal plume on aspiration efficiency of micro-particles. *Building and Environment*, 118:159–172, June 2017. ISSN 03601323. doi: 10.1016/j.buildenv.2017.03.018.
- [53] Quidel. Sofia SARS Antigen FIA Instructions for Use. <https://www.fda.gov/media/137885/download>. Technical report.
- [54] Paddy Robertson. Comparison of mask standards, ratings, and filtration effectiveness, June 2020. <https://smartairfilters.com/en/blog/comparison-mask-standards-rating-effectiveness/>.
- [55] Rebecca E. Stockwell, Michelle E. Wood, Congrong He, Laura J. Sherard, Emma L. Ballard, Timothy J. Kidd, Graham R. Johnson, Luke D. Knibbs, Lidia Morawska, Scott C. Bell, Maureen Peasey, Christine Duplancic, Kay A. Ramsay, Nassib Jabbour, Peter O’Rourke, Claire E. Wainwright, and Peter D. Sly. Face Masks Reduce the Release of *Pseudomonas aeruginosa* Cough Aerosols When Worn for Clinically Relevant Periods. *American Journal of Respiratory and Critical Care Medicine*, 198(10):1339–1342, November 2018. ISSN 1073-449X, 1535-4970. doi: 10.1164/rccm.201805-0823LE.
- [56] Baoqing Sun, Ying Feng, Xiaoneng Mo, Peiyan Zheng, Qian Wang, Pingchao Li, Ping Peng, Xiaoqing Liu, Zhilong Chen, Huimin Huang, Fan Zhang, Wenting Luo, Xuefeng Niu, Peiyu Hu, Longyu Wang, Hui Peng, Zhifeng Huang, Liqiang Feng, Feng Li, Fuchun Zhang, Fang Li, Nanshan Zhong, and Ling Chen. Kinetics of SARS-CoV-2 specific IgM and IgG responses in COVID-19 patients. *Emerging Microbes & Infections*, 9(1):940–948, 2020. doi: 10.1080/22221751.2020.1762515. URL <https://doi.org/10.1080/22221751.2020.1762515>. PMID: 32357808.
- [57] Julian W. Tang, Thomas J. Liebner, Brent A. Craven, and Gary S. Settles. A schlieren optical study of the human cough with and without wearing masks for aerosol infection control. *Journal of The Royal Society Interface*, 6(suppl_6), December 2009. ISSN 1742-5689, 1742-5662. doi: 10.1098/rsif.2009.0295.focus.
- [58] R. Tellier et al. Recognition of aerosol transmission of infectious agents: a commentary. *BMC Infectious Diseases*, 19(101):1–9, 2019.

- [59] Kelvin Kai-Wang To, Owen Tak-Yin Tsang, Wai-Shing Leung, Anthony Raymond Tam, Tak-Chiu Wu, David Christopher Lung, Cyril Chik-Yan Yip, Jian-Piao Cai, Jacky Man-Chun Chan, Thomas Shiu-Hong Chik, Daphne Pui-Ling Lau, Chris Yau-Chung Choi, Lin-Lei Chen, Wan-Mui Chan, Kwok-Hung Chan, Jonathan Daniel Ip, Anthony Chin-Ki Ng, Rosana Wing-Shan Poon, Cui-Ting Luo, Vincent Chi-Chung Cheng, Jasper Fuk-Woo Chan, Ivan Fan-Ngai Hung, Zhiwei Chen, Honglin Chen, and Kwok-Yung Yuen. Temporal profiles of viral load in posterior oropharyngeal saliva samples and serum antibody responses during infection by SARS-CoV-2: an observational cohort study. *The Lancet Infectious Diseases*, 20(5):565 – 574, 2020. ISSN 1473-3099. doi: [https://doi.org/10.1016/S1473-3099\(20\)30196-1](https://doi.org/10.1016/S1473-3099(20)30196-1). URL <http://www.sciencedirect.com/science/article/pii/S1473309920301961>.
- [60] Neeltje van Doremalen, Trenton Bushmaker, Dylan H. Morris, Myndi G. Holbrook, Amandine Gamble, Brandi N. Williamson, Azaibi Tamin, Jennifer L. Harcourt, Natalie J. Thornburg, Susan I. Gerber, James O. Lloyd-Smith, Emmie de Wit, and Vincent J. Munster. Aerosol and Surface Stability of SARS-CoV-2 as Compared with SARS-CoV-1. *New England Journal of Medicine*, page NEJMc2004973, March 2020. ISSN 0028-4793, 1533-4406. doi: 10.1056/NEJMc2004973.
- [61] Chantal B.F. Vogels, Anderson F. Brito, Anne Louise Wyllie, Joseph R Fauver, Isabel M. Ott, Chaney C. Kalinich, Mary E. Petrone, Arnau Casanovas-Massana, M. Catherine Muenker, Adam J. Moore, Jonathan Klein, Peiwen Lu, Alice Lu-Culligan, Xiaodong Jiang, Daniel J. Kim, Eriko Kudo, Tianyang Mao, Miyu Moriyama, Ji Eun Oh, Annsea Park, Julio Silva, Eric Song, Takehiro Takehashi, Manabu Taura, Maria Tokuyama, Arvind Venkataraman, Orr-El Weizman, Patrick Wong, Yexin Yang, Nagarjuna R. Cheemarla, Elizabeth White, Sarah Lapidus, Rebecca Earnest, Bertie Geng, Pavithra Vijayakumar, Camila Odio, John Fournier, Santos Bermejo, Shelli Farhadian, Charles Dela Cruz, Akiko Iwasaki, Albert I. Ko, Marie-Louise Landry, Ellen F. Foxman, and Nathan D. Grubaugh. Analytical sensitivity and efficiency comparisons of SARS-COV-2 qRT-PCR assays. Preprint, Infectious Diseases (except HIV/AIDS), April 2020.
- [62] Ania Wajnberg, Mayce Mansour, Emily Leven, Nicole M Bouvier, Gopi Patel, Adolfo Firpo, Rao Mendu, Jeffrey Jhang, Suzanne Arinsburg, Melissa Gitman, Jane Houldsworth, Ian Baine, Viviana Simon, Judith Aberg, Florian Krammer, David Reich, and Carlos Cordon-Cardo. Humoral immune response and prolonged PCR positivity in a cohort of

- 1343 SARS-CoV 2 patients in the New York City region. *medRxiv*, 2020. doi: 10.1101/2020.04.30.20085613. URL <https://www.medrxiv.org/content/early/2020/05/05/2020.04.30.20085613>.
- [63] N. Wilson, A. Kvalsvig, L. Barnard, and M. G Baker. Case-Fatality Risk Estimates for COVID-19 Calculated by Using a Lag Time for Fatality. *Emerging Infectious Diseases*, 26(6), 2020. doi: 10.3201/eid2606.200320.
- [64] M. E. J. Woolhouse, C. Dye, J.-F. Etard, T. Smith, J. D. Charlwood, G. P. Garnett, P. Hagan, J. L. K. Hii, P. D. Ndhlovu, R. J. Quinnell, C. H. Watts, S. K. Chandiwana, and R. M. Anderson. Heterogeneities in the transmission of infectious agents: Implications for the design of control programs. *Proceedings of the National Academy of Sciences*, 94(1):338–342, 1997. ISSN 0027-8424. doi: 10.1073/pnas.94.1.338. URL <https://www.pnas.org/content/94/1/338>.
- [65] Li-Ping Wu, Nai-Chang Wang, Yi-Hua Chang, Xiang-Yi Tian, Dan-Yu Na, Li-Yuan Zhang, Lei Zheng, Tao Lan, Lin-Fa Wang, and Guo-Dong Liang. Duration of Antibody Responses after Severe Acute Respiratory Syndrome. *Emerging Infectious Diseases*, 13(10):1562–1564, October 2007. ISSN 1080-6040, 1080-6059. doi: 10.3201/eid1310.070576.
- [66] R. Wölfel, V.M. Corman, and W. et al Guggemos. Virological assessment of hospitalized patients with COVID-2019. *Nature*, 581, 2020. doi: 10.1038/s41586-020-2196-x.
- [67] X. Xie, Y. Li, A. T. Y. Chwang, P. L. Ho, and W. H. Seto. How far droplets can move in indoor environments ? revisiting the Wells evaporation?falling curve. *Indoor Air*, 17(3):211–225, June 2007. ISSN 0905-6947, 1600-0668. doi: 10.1111/j.1600-0668.2007.00469.x.
- [68] Xiaojian Xie, Yuguo Li, Hequan Sun, and Li Liu. Exhaled droplets due to talking and coughing. *Journal of The Royal Society Interface*, 6(suppl_6), December 2009. ISSN 1742-5689, 1742-5662. doi: 10.1098/rsif.2009.0388.focus.
- [69] R. Yan, S. Chillrud, D.L. Magadini, and B. Yan. Developing home-disinfection and filtration efficiency improvement methods for N95 respirators and surgical facial masks: stretching supplies and better protection during the ongoing COVID-19 pandemic. *Journal of the International Society for Respiratory Protection*, 37(1), 2020.
- [70] Juanjuan Zhao, Quan Yuan, Haiyan Wang, Wei Liu, Xuejiao Liao, Yingying Su, Xin Wang, Jing Yuan, Tingdong Li, Jinxiu Li, Shen Qian, Congming Hong, Fuxiang Wang, Yingxia Liu, Zhaoqin Wang, Qing He,

Zhiyong Li, Bin He, Tianying Zhang, Yang Fu, Shengxiang Ge, Lei Liu, Jun Zhang, Ningshao Xia, and Zheng Zhang. Antibody responses to SARS-CoV-2 in patients of novel coronavirus disease 2019. *Clinical Infectious Diseases*, 03 2020. ISSN 1058-4838. doi: 10.1093/cid/ciaa344. URL <https://doi.org/10.1093/cid/ciaa344>.