

Online Appendix to “An Upper Bound to the Benefits
of Implementing Positive Assortative Matching in
Pooled Testing”

Gustavo Quinderé Saraiva

January 13, 2026

Contents

A Proofs	2
A.1 Expected Number of Tests	2
A.2 Expected Number of False Positives	11
A.3 Expected Number of False Negatives	22
A.4 Upper bounds with partial information on μ , a , b , S_e and S_p	23
B How large N needs to be for the upper bounds to be accurate?	24
C Alternative Parametrization for the Chlamydia Case Study	27
D Upper bounds without information on the dilution effect	29
D.1 Upper bound to the Expected Number of Tests without information regarding the dilution effect	30
D.2 Upper bound to the Expected Number of False Positives without information regarding the dilution effect	33
D.3 Upper bound to the Expected Number of False Negatives with partial information regarding the dilution effect	35
D.4 Numerical exercises with looser upper bounds	38
Bibliography	41

A Proofs

A.1 Expected Number of Tests

To avoid clutter notation, in this and the following sections we will sometimes redefine:

$$m_n(\bar{q}) = m_{n,K}(\bar{q}), \quad m_n = m_{n,K}(\mu),$$

$$c_n(\bar{q}) = c_{n,K}(\bar{q}), \quad c_n = c_{n,K}(\mu).$$

Lemma A.1 (*Online Appendix from Saraiva [2023]*) *If $h(\cdot, K)$ is concave, then, for any arbitrary group $G_g \subseteq S$ such that $|G_g| = K$, and any $l \in G_g$,*

$$\sum_{I=0}^{K-1} P_{G_g \setminus \{l\}}(I)[h(I+1, K) - h(I, K)]$$

is decreasing in the probability of infection from each subject in $G_g \setminus \{l\}$.

Lemma A.2 *Suppose that the order statistics $(q_1, q_2, \dots, q_N)$ and $(\hat{q}_1, \hat{q}_2, \dots, \hat{q}_N)$ are such that:*

1. $q_i, \hat{q}_i \in [a, b]$ for all i ,
2. There is no $i < j$ with $i \in G_g$ and $j \in G_{g'}$ where $g < g'$ such that $a < q_i \leq q_j < b$ or $a < \hat{q}_i \leq \hat{q}_j < b$, (i.e., there can be no more than one group in which the probability of infection from every subject within the group does not match the upper or the lower bound),
3. There is no i, j with $i, j \in G_g$ such that $q_i < q_j$ or $\hat{q}_i < \hat{q}_j$ (i.e., subjects belonging to the same group must have the same probability of infection),
4. And suppose that

$$\sum_{i=1}^N \frac{q_i}{N} = \sum_{i=1}^N \frac{\hat{q}_i}{N} = \bar{q}.$$

Then, $(q_1, q_2, \dots, q_N) = (\hat{q}_1, \hat{q}_2, \dots, \hat{q}_N) = q(\bar{q}, n, K)$.

Proof: Define

$$\begin{aligned} n_a &\equiv \#\{i; q_i = a\}, & n_b &\equiv \#\{i; q_i = b\} \\ \hat{n}_a &\equiv \#\{i; \hat{q}_i = a\}, & \hat{n}_b &\equiv \#\{i; \hat{q}_i = b\}. \end{aligned}$$

Then we must have $n_a + n_b \geq N - K$ and $\hat{n}_a + \hat{n}_b \geq N - K$, which implies that there is a $c \in (a, b]$ and a $\hat{c} \in (a, b]$ such that

$$\sum_i q_i = n_a a + (N - n_a - K)b + Kc$$

and

$$\sum_i \hat{q}_i = \hat{n}_a a + (N - \hat{n}_a - K)b + K\hat{c}.$$

Suppose by way of contradiction that $(q_1, q_2, \dots, q_N) \neq (\hat{q}_1, \hat{q}_2, \dots, \hat{q}_N)$. This happens if and only if $n_a \neq \hat{n}_a$. Without loss of generality suppose that $n_a < \hat{n}_a$, then $\hat{n}_a - n_a \geq K$.

Then, we must have

$$\begin{aligned}
\underbrace{\sum_i q_i - \sum_i \hat{q}_i}_{=0} &= n_a a + (N - n_a - 1)b + Kc - (\hat{n}_a a + (N - \hat{n}_a - 1)b + K\hat{c}) \\
\iff 0 &= (b - a)(\hat{n}_a - n_a) + K(c - \hat{c}) \\
\iff K(\hat{c} - c) &= (b - a) \underbrace{(\hat{n}_a - n_a)}_{\geq K} \\
\Rightarrow \hat{c} - c &\geq b - a,
\end{aligned}$$

a contradiction with $c \in (a, b]$ and a $\hat{c} \in (a, b]$. The rest of the proof follows from the fact that $(q_1, q_2, \dots, q_N)$ such that:

$$q_i = \begin{cases} a, & \text{if } i \leq n_a \\ \frac{N\bar{q} - n_a a - (N - n_a - K)b}{K}, & \text{if } n_a < i \leq n_a + K \\ b, & \text{if } i > n_a + K \end{cases} \quad \forall i$$

satisfy properties 1 to 4, so, from our previous result, this must be the only realization of order statistics that satisfy these properties. ■

Proof of proposition 2: Suppose that the order statistics $(q_1, \dots, q_N)$ minimize $T_o^K(q_1, \dots, q_N)$ subject to the constraints $\sum_i q_i/N = \bar{q}$ and $a \leq q_1 \leq q_N \leq b$.

Suppose by way of contradiction that there is a q_i and q_j with $i \neq j$ and $s_i \in G_g$ and $s_j \in G_{g'}$ with $g \neq g'$ (i.e., the subjects with probability of infection q_i and q_j belong to different groups) such that $a < q_i \leq q_j < b$. Then there is an $\varepsilon > 0$ such that $q_i - \varepsilon \geq a$ and $q_j + \varepsilon \leq b$. So replacing q_i by $q_i - \varepsilon$ and q_j by $q_j + \varepsilon$ does not affect the average probability of infection, and the restrictions $a \leq q_i \leq b \forall i$ still hold.

Let t_g be the expected number of tests from group G_g and let $t_{g'}$ be the expected number of tests from group $G_{g'}$. Let $\tilde{t}_g$ be the expected number of tests from group G_g after the probability of infection q_i from subject s_i is replaced by $q_i - \varepsilon$, and let $\tilde{t}_{g'}$ be the expected number of tests from group $G_{g'}$ after the probability of infection q_j from subject s_j is replaced by $q_j + \varepsilon$. We want to show that $\tilde{t}_g + \tilde{t}_{g'} \leq t_g + t_{g'}$. Now notice that

$$\begin{aligned}
&\tilde{t}_g + \tilde{t}_{g'} \leq t_g + t_{g'} \\
\iff &(q_i - \varepsilon) \sum_{I=0}^{K-1} P_{G_g \setminus \{i\}}(I) h(I+1) + (1 - q_i + \varepsilon) \sum_{I=0}^{K-1} P_{G_g \setminus \{i\}}(I) h(I) + \\
&+ (q_j + \varepsilon) \sum_{I=0}^{K-1} P_{G_{g'} \setminus \{j\}}(I) h(I+1) + (1 - q_j - \varepsilon) \sum_{I=0}^{K-1} P_{G_{g'} \setminus \{j\}}(I) h(I) \leq
\end{aligned}$$

$$\begin{aligned}
& q_i \sum_{I=0}^{K-1} P_{G_g \setminus \{i\}}(I) h(I+1) + (1 - q_i) \sum_{I=0}^{K-1} P_{G_g \setminus \{i\}}(I) h(I) + \\
& + q_j \sum_{I=0}^{K-1} P_{G_{g'} \setminus \{j\}}(I) h(I+1) + (1 - q_j) \sum_{I=0}^{K-1} P_{G_{g'} \setminus \{j\}}(I) h(I) \\
& \iff \varepsilon \sum_{I=0}^{K-1} \left[P_{G_{g'} \setminus \{j\}}(I) h(I+1) - P_{G_g \setminus \{i\}}(I) h(I) \right] \\
& \leq \varepsilon \sum_{I=0}^{K-1} \left[P_{G_g \setminus \{i\}}(I) h(I+1) - P_{G_g \setminus \{i\}}(I) h(I) \right] \tag{1}
\end{aligned}$$

Because the probability of infection from each subject in $G_g \setminus \{i\}$ is lower than the probability of infection from each subject in $G_{g'} \setminus \{j\}$ and because these groups have the same size, it then follows from lemma A.1 that inequality 1 holds.

Now, suppose by way of contradiction that there is a subject s_i (with probability of infection q_i) and a subject s_j (with probability of infection q_j) who both belong to the same group G_g where $a \leq q_i < q_j \leq b$. Then, for a given $\varepsilon \in [0, \frac{q_j - q_i}{2}]$ we have that $q_i + \varepsilon \in [a, q_j]$ and $q_j - \varepsilon \in [q_i, b]$. Let t_g be the expected number of tests from group G_g and let t_g^ε be the expected number of tests from group G_g after the probability of infection q_i from subject s_i is replaced by $q_i + \varepsilon$ and the probability of infection q_j from subject s_j is replaced by $q_j - \varepsilon$. We want to show that $t_g^\varepsilon \leq t_g$. Notice that

$$\begin{aligned}
& t_g^\varepsilon \leq t_g \\
& \iff (q_i + \varepsilon)(q_j - \varepsilon) \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I) h(I+2) \\
& + [(q_i + \varepsilon)(1 - q_j + \varepsilon) + (q_j - \varepsilon)(1 - q_i - \varepsilon)] \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I) h(I+1) \\
& + (1 - q_i - \varepsilon)(1 - q_j + \varepsilon) \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I) h(I) \leq \\
& q_i q_j \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I) h(I+2) + [q_i(1 - q_j) + q_j(1 - q_i)] \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I) h(I+1) \\
& + (1 - q_i)(1 - q_j) \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I) h(I) \\
& \iff \underbrace{\varepsilon(q_j - q_i - \varepsilon)}_{>0} \left[\left(\sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I) h(I+2) - \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I) h(I+1) \right) \right. \\
& \left. - \left(\sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I) h(I+1) - \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I) h(I) \right) \right] \leq 0
\end{aligned}$$

$$\Longleftrightarrow \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I) [h(I+2) - h(I+1)] \leq \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I) [h(I+1) - h(I)],$$

a condition that holds since the concavity of $h(\cdot)$ implies that $h(I+1) - h(I)$ is decreasing in I .

Therefore, $(q_1, q_2, \dots, q_N)$ must satisfy conditions 1 to 4 from lemma A.2, which implies that $(q_1, q_2, \dots, q_N) = q(\bar{q}, n, K)$. $\blacksquare$

Lemma A.3

$$\lim_{n \rightarrow \infty} \frac{m_{n,K}(\mu)}{N} = \frac{b - \mu}{b - a}.$$

Proof: Notice that, for all $n \in \mathbb{N}$, we have that

$$\mu = \frac{m_n}{N}a + \left(1 - \frac{m_n}{N} - \frac{K}{N}\right)b + \frac{c_{n,K}(\mu)K}{N},$$

where $c_{n,K}(\mu) \in [a, b]$. Because $\{c_{n,K}(\mu)\}_n$ is a bounded sequence, we have that $\lim_{n \rightarrow \infty} \frac{c_{n,K}(\mu)K}{N} = 0$. Therefore,

$$\begin{aligned} & \lim_{n \rightarrow \infty} \left[\frac{m_n}{N}a + \left(1 - \frac{m_n}{N} - \underbrace{\frac{K}{N}}_{\rightarrow 0}\right)b \right] = \mu \\ \Rightarrow & \lim_{n \rightarrow \infty} \left[\frac{m_n}{N}a + \left(1 - \frac{m_n}{N}\right)b \right] = \mu \\ \Rightarrow & \lim_{n \rightarrow \infty} \frac{m_n}{N} = \frac{b - \mu}{b - a}. \end{aligned}$$

$\blacksquare$

Lemma A.4

$$\frac{m_{n,K}(\bar{q})}{N} \xrightarrow{p} \frac{b - \mu}{b - a}.$$

Proof: Notice that, for all $n \in \mathbb{N}$ and every realization of $\bar{q}$, we have that

$$\bar{q} = \frac{m_n(\bar{q})}{N}a + \left(1 - \frac{m_n(\bar{q})}{N} - \frac{K}{N}\right)b + \frac{c_{n,K}(\bar{q})K}{N},$$

where $c_{n,K}(\bar{q}) \in [a, b]$. By the law of large numbers we know that

$$\begin{aligned} & \bar{q} \xrightarrow{p} \mu \\ \Longleftrightarrow & \frac{m_n(\bar{q})}{N}a + \left(1 - \frac{m_n(\bar{q})}{N} - \frac{K}{N}\right)b + \frac{c_{n,K}(\bar{q})K}{N} \xrightarrow{p} \mu \end{aligned}$$

Because the sequence $\{c_{n,K}(\bar{q})\}_n$ of random variables is bounded, we have that $\frac{c_{n,K}(\bar{q})K}{N} \xrightarrow{p} 0$. Therefore,

$$\begin{aligned} & \frac{m_n(\bar{q})}{N}a + \left(1 - \frac{m_n(\bar{q})}{N} - \underbrace{\frac{K}{N}}_{\rightarrow 0}\right)b \xrightarrow{p} \mu \\ \Rightarrow & \frac{m_n(\bar{q})}{N}a + \left(1 - \frac{m_n(\bar{q})}{N}\right)b \xrightarrow{p} \mu \\ \Rightarrow & \frac{m_n(\bar{q})}{N} \xrightarrow{p} \frac{b - \mu}{b - a}. \end{aligned}$$

■

Proof of theorem 1: For each $n \in \mathbb{N}$, we have that

$$\begin{aligned} UB_o^T(\mu, n, K) = & \sum_{I=0}^K h(I, K) \binom{K}{I} \left[\mu^I (1 - \mu)^{K-I} - \frac{m_n}{N} a^I (1 - a)^{K-I} \right. \\ & \left. - \left(1 - \frac{m_n}{N} - \frac{K}{N}\right) b^I (1 - b)^{K-I} - \frac{K}{N} c_n^I (1 - c_n)^{K-I} \right]. \end{aligned}$$

Clearly, $\lim_{n \rightarrow \infty} \frac{K}{N} \sum_{I=0}^K h(I, K) \binom{K}{I} c_n^I (1 - c_n)^{K-I} = 0$. Therefore,

$$\begin{aligned} \lim_{n \rightarrow \infty} UB_o^T(\mu, n, K) &= \lim_{n \rightarrow \infty} \sum_{I=0}^K h(I, K) \binom{K}{I} \left[\mu^I (1 - \mu)^{K-I} \right. \\ & \quad \left. - \frac{m_n}{N} a^I (1 - a)^{K-I} - \left(1 - \frac{m_n}{N} - \underbrace{\frac{K}{N}}_{\rightarrow 0}\right) b^I (1 - b)^{K-I} \right] \\ &= \lim_{n \rightarrow \infty} \sum_{I=0}^K h(I, K) \binom{K}{I} \left[\mu^I (1 - \mu)^{K-I} - \frac{m_n}{N} a^I (1 - a)^{K-I} \right. \\ & \quad \left. - \left(1 - \frac{m_n}{N}\right) b^I (1 - b)^{K-I} \right]. \end{aligned} \tag{2}$$

But from lemma A.3 we have that $\lim_{n \rightarrow \infty} \frac{m_n}{N} = \frac{b - \mu}{b - a}$, which, from expression (2) implies that

$$\begin{aligned} \lim_{n \rightarrow \infty} UB_o^T(\mu, n, K) &= \sum_{I=0}^K h(I, K) \binom{K}{I} \left[\mu^I (1 - \mu)^{K-I} \right. \\ & \quad \left. - \frac{b - \mu}{b - a} a^I (1 - a)^{K-I} - \left(1 - \frac{b - \mu}{b - a}\right) b^I (1 - b)^{K-I} \right], \end{aligned}$$

as we wanted to show.

To prove that

$$UB_o^T(\bar{q}, n, K) \xrightarrow{p} \sum_{I=0}^K h(I, K) \binom{K}{I} \left[\mu^I (1 - \mu)^{K-I} - \frac{b - \mu}{b - a} a^I (1 - a)^{K-I} - \left(1 - \frac{b - \mu}{b - a}\right) b^I (1 - b)^{K-I} \right],$$

we use an analogous argument, but instead of invoking lemma A.3 we invoke lemma A.4 and properties of convergence in probability. Indeed, notice that

$$UB_o^T(\bar{q}, n, K) = \sum_{I=0}^K h(I, K) \binom{K}{I} \left[\bar{q}^I (1 - \bar{q})^{K-I} - \frac{m_n(\bar{q})}{N} a^I (1 - a)^{K-I} - \left(1 - \frac{m_n(\bar{q})}{N} - \underbrace{\frac{K}{N}}_{\rightarrow 0}\right) b^I (1 - b)^{K-I} - \underbrace{\frac{K}{N} c_n(\bar{q})^I (1 - c_n(\bar{q}))^{K-I}}_{\xrightarrow{p} 0} \right].$$

Because the function $f(x) = x^I(1 - x)^{K-I}$ is continuous, it then follows from the law of large numbers and the continuous mapping theorem that $\bar{q}^I(1 - \bar{q})^{K-I} \xrightarrow{p} \mu^I(1 - \mu)^{K-I}$ for all $I \leq K$. Moreover, the sequence $\frac{K}{N} b^I(1 - b)^{K-I}$ converges to 0 as n goes to infinity, and $\frac{K}{N} c_n(\bar{q})^I(1 - c_n(\bar{q}))^{K-I} \xrightarrow{p} 0$. So, it suffices to show that the sequence of random variables $\frac{m_n(\bar{q})}{N}$ converges in probability to $\alpha = \frac{b - \mu}{b - a}$. From lemma A.4, this condition holds.

Finally, $T_r^K(\bar{q}, n, K) \xrightarrow{p} T_r^K(\mu, n, K)$ by the continuous mapping theorem. ■

Proof of proposition 3: Defining

$$\theta_n \equiv \frac{\alpha - m_n/N}{K/N}$$

and

$$c_n = \theta_n a + (1 - \theta_n) b,$$

we have that

$$UB_o^T(\mu, n, K) = \sum_{I=0}^K h(I, K) \binom{K}{I} \left[\mu^I (1 - \mu)^{K-I} - \frac{m_n}{N} a^I (1 - a)^{K-I} - \left(1 - \frac{m_n}{N} - \frac{K}{N}\right) b^I (1 - b)^{K-I} - \frac{K}{N} c_n^I (1 - c_n)^{K-I} \right].$$

So, to prove that

$$UB_o^T(\mu, n, K) \leq \sum_{I=0}^K h(I, K) \binom{K}{I} \left[\mu^I (1 - \mu)^{K-I} - \alpha a^I (1 - a)^{K-I} - (1 - \alpha) b^I (1 - b)^{K-I} \right], \quad (3)$$

it suffices to show that

$$\sum_{I=0}^K h(I, K) \binom{K}{I} [\theta_n a^I (1-a)^{K-I} + (1-\theta_n) b^I (1-b)^{K-I} - c_n^I (1-c_n)^{K-I}] \leq 0 \quad (4)$$

for all $n \in \mathbb{N}$.

To show that inequality (4) holds, consider two groups, G_g and $G_{g'}$, where the order statistics from group G_g are given by $(q_1, q_2, \dots, q_K)$ and the ones from group $G_{g'}$ are given by $(q'_1, q'_2, \dots, q'_K)$, where $q_i \leq q'_j \forall i, j$. With this notation, consider the following minimization problem:

$$\begin{aligned} & \min_{(q_1, \dots, q_K), (q'_1, \dots, q'_K)} \theta \sum_{I=0}^K P_{G_g}(I) h(I) + (1-\theta) \sum_{I=0}^K P_{G_{g'}}(I) h(I) \\ \text{s.t. } & \theta \frac{\sum_i q_i}{K} + (1-\theta) \frac{\sum_i \tilde{q}_i}{K} = c_n \\ & a \leq q_i \leq q'_j \leq b \quad \forall i, j \end{aligned} \quad (5)$$

So, if $(q_1, \dots, q_K) = (a, a, \dots, a)$ and $(q'_1, \dots, q'_K) = (b, b, \dots, b)$ is a solution to this minimization problem we have that inequality (4) holds.

Suppose by way of contradiction that $(q_1, \dots, q_K)$ and $(q'_1, \dots, q'_K)$ solve this minimization problem and there is a $q_i > a$ and a $q'_j < b$. Then, there is an $\varepsilon > 0$ such that $q_i - \frac{\varepsilon}{\theta} > a$ and $q'_j + \frac{\varepsilon}{1-\theta} < b$.

Let t_g be the expected number of tests from group G_g (i.e., $t_g = \sum_{I=0}^K P_{G_g}(I) h(I)$) and let $t_{g'}$ be the expected number of tests from group $G_{g'}$ (i.e., $\sum_{I=0}^K P_{G_{g'}}(I) h(I)$). Let $\tilde{t}_g$ be the expected number of tests from group G_g after the probability of infection q_i from subject s_i is replaced by $q_i - \frac{\varepsilon}{\theta}$, and let $\tilde{t}_{g'}$ be the expected number of tests from group $G_{g'}$ after the probability of infection q_j from subject s_j is replaced by $q_j + \frac{\varepsilon}{1-\theta}$. We want to show that $\theta \tilde{t}_g + (1-\theta) \tilde{t}_{g'} \leq \theta t_g + (1-\theta) t_{g'}$. Now notice that

$$\begin{aligned} & \theta \tilde{t}_g + (1-\theta) \tilde{t}_{g'} \leq \theta t_g + (1-\theta) t_{g'} \\ \iff & \theta \left(q_i - \frac{\varepsilon}{\theta} \right) \sum_{I=0}^{K-1} P_{G_g \setminus \{i\}}(I) h(I+1) + \theta \left(1 - q_i + \frac{\varepsilon}{\theta} \right) \sum_{I=0}^{K-1} P_{G_g \setminus \{i\}}(I) h(I) + \\ & + (1-\theta) \left(q'_j + \frac{\varepsilon}{1-\theta} \right) \sum_{I=0}^{K-1} P_{G_{g'} \setminus \{j\}}(I) h(I+1) \\ & + (1-\theta) \left(1 - q'_j - \frac{\varepsilon}{1-\theta} \right) \sum_{I=0}^{K-1} P_{G_{g'} \setminus \{j\}}(I) h(I) \leq \\ & \theta q_i \sum_{I=0}^{K-1} P_{G_g \setminus \{i\}}(I) h(I+1) + \theta (1 - q_i) \sum_{I=0}^{K-1} P_{G_g \setminus \{i\}}(I) h(I) + \\ & (1 - q'_j) \sum_{I=0}^{K-1} P_{G_{g'} \setminus \{j\}}(I) h(I+1) + q'_j \sum_{I=0}^{K-1} P_{G_{g'} \setminus \{j\}}(I) h(I) \end{aligned}$$

$$\begin{aligned}
& + (1 - \theta)q_j \sum_{I=0}^{K-1} P_{G_{g'} \setminus \{j\}}(I)h(I+1) + (1 - \theta)(1 - q_j) \sum_{I=0}^{K-1} P_{G_{g'} \setminus \{i\}}(I)h(I) \\
& \iff \varepsilon \sum_{I=0}^{K-1} \left[P_{G_{g'} \setminus \{j\}}(I)h(I+1) - P_{G_{g'} \setminus \{j\}}(I)h(I) \right] \leq \\
& \varepsilon \sum_{I=0}^{K-1} \left[P_{G_g \setminus \{i\}}(I)h(I+1) - P_{G_g \setminus \{i\}}(I)h(I) \right]
\end{aligned} \tag{6}$$

Because the probability of infection from each subject in $G_g \setminus \{i\}$ is lower than the probability of infection from each subject in $G_{g'} \setminus \{j\}$ and because these groups have the same size, it then follows from lemma A.1 that inequality (6) holds.

Therefore, if $(q_1, \dots, q_K)$ and $(\tilde{q}_1, \dots, \tilde{q}_K)$ is a solution to the minimization problem (5), we must either have $(q_1, \dots, q_K) = (a, a, \dots, a)$, which implies that $(q'_1, \dots, q'_K) = (b, b, \dots, b)$, or we must have $(q'_1, \dots, q'_K) = (b, b, \dots, b)$, which implies that $(q_1, \dots, q_K) = (a, a, \dots, a)$. So $(q_1, \dots, q_K) = (a, a, \dots, a)$ and $(\tilde{q}_1, \dots, \tilde{q}_K) = (b, b, \dots, b)$ is the only candidate to solve this minimization problem.

Clearly, the feasible set from this minimization problem is non-empty, as $(q_1, \dots, q_K) = (a, a, \dots, a)$ and $(q'_1, \dots, q'_K) = (b, b, \dots, b)$ satisfy all the constraints from this minimization problem. Moreover, the feasible set is compact, which implies that a solution to the minimization problem (5) exists. So, because $((q_1, \dots, q_K), (q'_1, \dots, q'_K)) = ((a, a, \dots, a), (b, b, \dots, b))$ is our only candidate for a solution, it must be the unique solution, as we wanted to show. ■

Proof of Proposition 4: That

$$UB_o^T(\mu, n, K) = (S_e + S_p - 1) [\alpha(1 - a)^K + (1 - \alpha)(1 - b)^K - (1 - \mu)^K]$$

in the absence of dilution effects can be easily derived from equation (5). The rest of the proof follows from the fact that $(S_e + S_p - 1) > 0$ and that

$$\alpha(1 - a)^K + (1 - \alpha)(1 - b)^K - (1 - \mu)^K = \frac{b - \mu}{b - a}(1 - a)^K + \frac{\mu - a}{b - a}(1 - b)^K - (1 - \mu)^K$$

is concave with respect to μ , which, from the Jensen's inequality, implies that

$$\alpha(1 - a)^K + (1 - \alpha)(1 - b)^K - (1 - \mu)^K \geq \mathbb{E}_{\bar{q}} \left[\frac{b - \bar{q}}{b - a}(1 - a)^K + \frac{\bar{q} - a}{b - a}(1 - b)^K - (1 - \mu)^K \right].$$

■

A.2 Expected Number of False Positives

For any group $G_g \subseteq S$ and any $j \in G_g$ define

$$C_{G_g,j} \equiv \sum_{I=0}^{K-1} P_{G_g \setminus \{j\}}(I) h(I+1)(K - (I+1))(1 - S_p)$$

and

$$D_{G_g,j} \equiv \sum_{I=0}^{K-1} P_{G_g \setminus \{j\}}(I) h(I)(K - I)(1 - S_p).$$

From the above expressions, $C_{G_g,j}$ is the expected number of false positives in group G_g conditional that subject $s_j \in G_g$ is infected. Similarly, $D_{G_g,j}$ is the expected number of false positives in group G_g conditional that subject $s_j \in G_g$ is not infected.

Now notice that, for any $j \in G_g$,

$$fp_g = q_j C_{G_g,j} + (1 - q_j) D_{G_g,j} = q_j (C_{G_g,j} - D_{G_g,j}) + D_{G_g,j}$$

Lemma A.5 (*Online Appendix from Saraiva [2023]*) *If $h(\cdot, K)$ is concave and $l \in G_g$, then*

$$C_{G_g,l} - D_{G_g,l}$$

is decreasing in the probability of infection from each subject in $G_g \setminus \{l\}$.

Proof of proposition 5: Suppose that the order statistics $(q_1, \dots, q_N)$ minimize $FP_o^K(q_1, \dots, q_N)$ subject to the constraints $\sum_i q_i / N = \bar{q}$ and $a \leq q_1 \leq q_N \leq b$.

Suppose by way of contradiction that there is a q_i and q_j with $i \neq j$ and $s_i \in G_g$ and $s_j \in G_{g'}$ with $g \neq g'$ (i.e., the subjects with probability of infection q_i and q_j belong to different groups) such that $a < q_i \leq q_j < b$. Then there is an $\varepsilon > 0$ such that $q_i - \varepsilon \geq a$ and $q_j + \varepsilon \leq b$. So replacing q_i by $q_i - \varepsilon$ and q_j by $q_j + \varepsilon$ does not affect the average probability of infection, and the restrictions $a \leq q_i \leq b \forall i$ still hold.

Let fp_g be the expected number of false positives from group G_g and let $fp_{g'}$ be the expected number of false positives from group $G_{g'}$. Let $\widetilde{fp}_g$ be the expected number of false positives from group G_g after the probability of infection q_i from subject s_i is replaced by $q_i - \varepsilon$, and let $\widetilde{fp}_{g'}$ be the expected number of false positives from group $G_{g'}$ after the probability of infection q_j from subject s_j is replaced by $q_j + \varepsilon$. We want to

show that $\widetilde{fp}_g + \widetilde{fp}_{g'} \leq fp_g + fp_{g'}$. Now notice that

$$\begin{aligned}
& \widetilde{fp}_g + \widetilde{fp}_{g'} \leq fp_g + fp_{g'} \\
& \iff (q_i - \varepsilon)(C_{G_g,i} - D_{G_g,i}) + (q_j + \varepsilon)(C_{G_{g'},i} - D_{G_{g'},j}) \leq \\
& \quad q_i(C_{G_g,i} - D_{G_g,i}) + q_j(C_{G_{g'},j} - D_{G_{g'},j}) \\
& \iff (C_{G_{g'},j} - D_{G_{g'},j}) \leq (C_{G_g,i} - D_{G_g,i}). \tag{7}
\end{aligned}$$

Because the probability of infection from each subject in $G_g \setminus \{i\}$ is lower than the probability of infection from each subject in $G_{g'} \setminus \{j\}$ and because these groups have the same size, it then follows from lemma A.1 that inequality (7) holds. Therefore, there can be no q_i and q_j with $i \neq j$ and $s_i \in G_g$ and $s_j \in G_{g'}$ with $g \neq g'$ such that $a < q_i \leq q_j < b$.

Now, suppose by way of contradiction that there is a $i, j \in G_g$ (i.e., subjects s_i and s_j belong to the same group G_g) where $a \leq q_i < q_j \leq b$. Then, for a given $\varepsilon \in [0, \frac{q_j - q_i}{2}]$ we have that $q_i + \varepsilon \in [a, q_j]$ and $q_j - \varepsilon \in [q_i, b]$. Let fp_g be the expected number of false positives from group G_g and let fp_g^ε be the expected number of false positives from group G_g after the probability of infection q_i from subject s_i is replaced by $q_i + \varepsilon$ and the probability of infection q_j from subject s_j is replaced by $q_j - \varepsilon$. We want to show that

$fp_g^\varepsilon \leq fp_g$. Notice that

$$\begin{aligned}
& fp_g^\varepsilon \leq fp_g \\
& \iff (q_i + \varepsilon)(q_j - \varepsilon) \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I)(K - I - 2)h(I + 2) \\
& \quad + [(q_i + \varepsilon)(1 - q_j + \varepsilon) + (q_j - \varepsilon)(1 - q_i - \varepsilon)] \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I)(K - I - 1)h(I + 1) \\
& \quad + (1 - q_i - \varepsilon)(1 - q_j + \varepsilon) \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I)(K - I)h(I) \leq \\
& \quad q_i q_j \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I)(K - I - 2)h(I + 2) \\
& \quad + [q_i(1 - q_j) + q_j(1 - q_i)] \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I)(K - I - 1)h(I + 1) \\
& \quad + (1 - q_i)(1 - q_j) \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I)(K - I)h(I) \\
& \iff \varepsilon \underbrace{(q_j - q_i - \varepsilon)}_{>0} \left[\left(\sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I)(K - I - 2)h(I + 2) \right. \right. \\
& \quad \left. \left. - \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I)(K - I - 1)h(I + 1) \right) \right. \\
& \quad \left. - \left(\sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I)(K - I - 1)h(I + 1) - \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I)(K - I)h(I) \right) \right] \leq 0 \\
& \iff \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I) [(K - I - 2)h(I + 2) - (K - I - 1)h(I + 1)] \leq \\
& \quad \sum_{I=0}^{K-2} P_{G_g \setminus \{i,j\}}(I) [(K - I - 1)h(I + 1) - (K - I)h(I)]. \tag{8}
\end{aligned}$$

Because $h(\cdot)$ is increasing and concave, we have that $[(K - I - 1)h(I + 1) - (K - I)h(I)]$ is decreasing in I , which implies that inequality (8) holds, from which we conclude that we cannot have $a \leq q_i < q_j \leq b$ (where $i, j \in G_g$).

Therefore, $(q_1, q_2, \dots, q_N)$ must satisfy conditions 1 to 4 from lemma A.2, which implies that $(q_1, q_2, \dots, q_N) = q(\bar{q}, n, K)$. ■

Proof of theorem 2: Noticing that:

$$UB_o^{FP}(\bar{q}, n, K) = (1 - S_p) \sum_{I=0}^{K-1} h(I, K) \binom{K-1}{I} \left[\mu^I (1 - \mu)^{K-I} - \frac{m_{n,K}(\bar{q})}{N} a^I (1 - a)^{K-I} \right. \\ \left. - (1 - \frac{m_{n,K}(\bar{q})}{N} - \frac{K}{N}) b^I (1 - b)^{K-I} - \frac{K}{N} c_n^I (1 - c_n)^{K-I} \right],$$

we can apply the same steps as the ones presented in the proof of theorem 1 to get the desired result. $\blacksquare$

Lemma A.6 For each $\mu \in [a, b]$ denote G_μ as a group with $K - 1$ subjects, each of which has probability of infection μ . Also, let

$$C_{G_\mu} \equiv (1 - S_p) \sum_{I=0}^{K-1} \binom{K-1}{I} \mu^I (1 - \mu)^{K-1-I} h(I+1) [K - (I+1)] \\ D_{G_\mu} \equiv (1 - S_p) \sum_{I=0}^{K-1} \binom{K-1}{I} \mu^I (1 - \mu)^{K-1-I} h(I) (K - I)$$

i.e., C_{G_μ} is the expected number of false positives from group G_μ conditional that one of its subjects is infected, whereas D_{G_μ} is the expected number of false positives from group G_μ conditional that one of its subjects is not infected.

Then, we have that

$$N(1 - S_p) \sum_{I=0}^{K-1} h(I) \binom{K-1}{I} \mu^I (1 - \mu)^{K-I} = n [\mu C_{G_\mu} + (1 - \mu) D_{G_\mu}].$$

Proof: Notice that

$$n [\mu C_{G_\mu} + (1 - \mu) D_{G_\mu}] = n(1 - S_p) \left[\sum_{I=0}^{K-1} \binom{K-1}{I} \mu^{I+1} (1 - \mu)^{K-1-I} h(I+1) [K - (I+1)] \right. \\ \left. + \sum_{I=0}^{K-1} \binom{K-1}{I} \mu^I (1 - \mu)^{K-I} h(I) (K - I) \right] \\ = n(1 - S_p) [(1 - \mu)^{K-1} h(0) K \\ + \left[\binom{K-1}{0} + \binom{K-1}{1} \right] \mu^1 (1 - \mu)^{K-1} h(1) (K - 1) \\ + \left[\binom{K-1}{1} + \binom{K-1}{2} \right] \mu^2 (1 - \mu)^{K-2} h(2) (K - 2) + \dots \\ + \left[\binom{K-1}{K-2} + \binom{K-1}{K-1} \right] \mu^{K-1} (1 - \mu) h(K - 1)] \quad (9)$$

Now, noticing that, for all $I \geq 1$

$$K \binom{K-1}{I} = (K - I) \left[\binom{K-1}{I-1} + \binom{K-1}{I} \right],$$

we can rewrite expression (9) as

$$\begin{aligned} n [\mu C_{G_\mu} + (1 - \mu) D_{G_\mu}] &= \overbrace{n}^{N/K} K(1 - S_p) \left[\sum_{I=0}^{K-1} \binom{K-1}{I} \mu^I (1 - \mu)^{K-I} h(I) \right] \\ &= N(1 - S_p) \sum_{i=0}^{K-1} \binom{K-1}{I} \mu^I (1 - \mu)^{K-I} h(I), \end{aligned}$$

as we wanted to show. ■

Proof of proposition 6: Defining

$$\theta_n \equiv \frac{\alpha - m_n/N}{K/N}$$

and

$$c_n = \theta_n a + (1 - \theta_n) b,$$

we have that

$$\begin{aligned} UB_o^{FP}(\mu, n, K) &= (1 - S_p) \sum_{I=0}^{K-1} h(I, K) \binom{K-1}{I} \left[\mu^I (1 - \mu)^{K-I} - \frac{m_n}{N} a^I (1 - a)^{K-I} \right. \\ &\quad \left. - (1 - \frac{m_n}{N} - \frac{K}{N}) b^I (1 - b)^{K-I} - \frac{K}{N} c_n^I (1 - c_n)^{K-I} \right]. \end{aligned} \quad (10)$$

So, to prove that

$$\begin{aligned} UB_o^{FP}(\mu, n, K) &\leq (1 - S_p) \sum_{I=0}^{K-1} h(I, K) \binom{K-1}{I} \left[\mu^I (1 - \mu)^{K-I} \right. \\ &\quad \left. - \alpha a^I (1 - a)^{K-I} - (1 - \alpha) b^I (1 - b)^{K-I} \right] \end{aligned} \quad (11)$$

it suffices to show that

$$\sum_{I=0}^K h(I, K) \binom{K-1}{I} [\theta_n a^I (1 - a)^{K-I} + (1 - \theta_n) b^I (1 - b)^{K-I} - c_n^I (1 - c_n)^{K-I}] \leq 0 \quad (12)$$

for all $n \in \mathbb{N}$.

From Lemma A.6, we can rewrite inequality (12) as

$$\theta_n (a C_{G_a} + (1 - a) D_{G_a}) + (1 - \theta_n) (b C_{G_b} + (1 - b) D_{G_b}) \leq c_n C_{G_{c_n}} + (1 - c_n) D_{G_{c_n}}, \quad (13)$$

where, for each $c \in [a, b]$, G_c corresponds to a group with $K - 1$ subjects, each with probability of infection c , and

$$C_{G_c} \equiv \sum_{I=0}^{K-1} P_{G_c}(I) h(I + 1) (K - (I + 1)) (1 - S_p),$$

$$D_{G_c} \equiv \sum_{I=0}^{K-1} P_{G_c}(I)h(I)(K-I)(1-S_p),$$

so that $cC_{G_c} + (1-c)D_{G_c}$ corresponds to the expected number of false positives from a group with K subjects, each with probability of infection c .

To show that inequality (13) holds, consider two groups, G_g and $G_{g'}$, where the order statistics from group G_g are given by $(q_1, q_2, \dots, q_K)$ and the ones from group $G_{g'}$ are given by $(q'_1, q'_2, \dots, q'_K)$, with $q_i \leq q'_j \forall i, j$. Let fp_g be the expected number of false positives from group G_g (i.e., for any $i \in \{1, 2, \dots, K\}$ $fp_g = q_i C_{G_g, i} + (1 - q_i) D_{G_g, i}$) and let $fp_{g'}$ be the expected number of false positives from group $G_{g'}$ (i.e., for any $i \in \{1, 2, \dots, K\}$ $fp_{g'} = q'_i C_{G_{g'}, i} + (1 - q'_i) D_{G_{g'}, i}$).

With this notation, consider the following minimization problem:

$$\begin{aligned} & \min_{(q_1, \dots, q_K), (q'_1, \dots, q'_K)} \theta fp_g + (1 - \theta) fp_{g'} \\ \text{s.t. } & \theta \frac{\sum_i q_i}{K} + (1 - \theta) \frac{\sum_i q'_i}{K} = c_n \\ & a \leq q_i \leq q'_j \leq b \quad \forall i, j \end{aligned} \tag{14}$$

If $(q_1, \dots, q_K) = (a, a, \dots, a)$ and $(q'_1, \dots, q'_K) = (b, b, \dots, b)$ is a solution to this minimization problem we have that inequality (12) holds.

Suppose by way of contradiction that $(q_1, \dots, q_K)$ and $(q'_1, \dots, q'_K)$ solve this minimization problem and there is a $q_i > a$ and a $q'_j < b$. Then, there is an $\varepsilon > 0$ such that $q_i - \frac{\varepsilon}{\theta} > a$ and $q'_j + \frac{\varepsilon}{1-\theta} < b$.

Let $\widetilde{fp}_g$ be the expected number of false positives from group G_g after the probability of infection q_i from subject s_i is replaced by $q_i - \frac{\varepsilon}{\theta}$, and let $\widetilde{fp}_{g'}$ be the expected number of false positives from group $G_{g'}$ after the probability of infection q_j from subject s_j is replaced by $q_j + \frac{\varepsilon}{1-\theta}$. We want to show that $\theta \widetilde{fp}_g + (1 - \theta) \widetilde{fp}_{g'} \leq \theta fp_g + (1 - \theta) fp_{g'}$. Now notice that

$$\begin{aligned} & \theta \widetilde{fp}_g + (1 - \theta) \widetilde{fp}_{g'} \leq \theta fp_g + (1 - \theta) fp_{g'} \\ \iff & \theta \left(q_i - \frac{\varepsilon}{\theta} \right) (C_{G_g, i} - D_{G_g, i}) + (1 - \theta) \left(q_j + \frac{\varepsilon}{1-\theta} \right) (C_{G_{g'}, i} - D_{G_{g'}, i}) \leq \\ & \theta q_i (C_{G_g, i} - D_{G_g, i}) + (1 - \theta) q_j (C_{G_{g'}, i} - D_{G_{g'}, i}) \\ \iff & (C_{G_{g'}, i} - D_{G_{g'}, i}) \leq (C_{G_g, i} - D_{G_g, i}) \end{aligned} \tag{15}$$

Because the probability of infection from each subject in $G_g \setminus \{i\}$ is lower than the

probability of infection from each subject in $G_{g'} \setminus \{j\}$ and because these groups have the same size, it then follows from Lemma A.5 that inequality (15) holds.

Therefore, if $(q_1, \dots, q_K)$ and $(\tilde{q}_1, \dots, \tilde{q}_K)$ is a solution to the minimization problem (14), we must either have $(q_1, \dots, q_K) = (a, a, \dots, a)$, which implies that $(q'_1, \dots, q'_K) = (b, b, \dots, b)$, or we must have $(q'_1, \dots, q'_K) = (b, b, \dots, b)$, which implies that $(q_1, \dots, q_K) = (a, a, \dots, a)$. So $(q_1, \dots, q_K) = (a, a, \dots, a)$ and $(\tilde{q}_1, \dots, \tilde{q}_K) = (b, b, \dots, b)$ is the only candidate to solve this minimization problem.

Clearly, the feasible set from this minimization problem is non-empty and compact, which implies that a solution to the minimization problem (14) exists. Therefore, because $((q_1, \dots, q_K), (q'_1, \dots, q'_K)) = ((a, a, \dots, a), (b, b, \dots, b))$ is our only candidate for a solution, it must be the unique solution, as we wanted to show. ■

Proof of Proposition 7: That

$$UB_o^{FP}(\mu, n, K) = [(1 - S_p)S_e - (1 - S_p)^2] [\alpha(1 - a)^K + (1 - \alpha)(1 - b)^K - (1 - \mu)^K]$$

in the absence of dilution effects can be easily derived from equation (9). The rest of the proof follows from the fact that $[(1 - S_p)S_e - (1 - S_p)^2] > 0$ and that

$$\alpha(1 - a)^K + (1 - \alpha)(1 - b)^K - (1 - \mu)^K = \frac{b - \mu}{b - a}(1 - a)^K + \frac{\mu - a}{b - a}(1 - b)^K - (1 - \mu)^K$$

is concave with respect to μ , which, from the Jensen's inequality, implies that

$$\alpha(1 - a)^K + (1 - \alpha)(1 - b)^K - (1 - \mu)^K \geq \mathbb{E}_{\bar{q}} \left[\frac{b - \bar{q}}{b - a}(1 - a)^K + \frac{\bar{q} - a}{b - a}(1 - b)^K - (1 - \mu)^K \right].$$

■

Proof of Proposition 12: Because N is even, matching all subjects into equal pools of size $K = 2$ is feasible.

Random Pooling: Suppose that random pooling is implemented.

Case 1: Suppose we had a pool with $K \geq 4$ subjects. In this case, the expected number of false positives from this pool would be given by

$$K [(1 - \mu)(1 - S_p)S_e - (1 - \mu)^K(1 - S_p)(S_e + S_p - 1)]. \quad (16)$$

If we took two subjects from this pool and matched them together to form a pool of size 2, the expected number of false positives from this new pool plus the expected

number of false positives from what remained of the original pool would be given by

$$2 \left[(1 - \mu)(1 - S_p)S_e - (1 - \mu)^2(1 - S_p)(S_e + S_p - 1) \right]. \quad (17)$$

Subtracting (16) by (17) yields:

$$\begin{aligned} & (1 - S_p)(S_e + S_p - 1) \left[(K - 2)(1 - \mu)^{K-2} + 2(1 - \mu)^2 - K(1 - \mu)^K \right] \geq \\ & (1 - S_p)(S_e + S_p - 1) \left[(K - 2)(1 - \mu)^{K-2} + 2(1 - \mu)^2 - K(1 - \mu)^{K-2} \right] = \\ & (1 - S_p)(S_e + S_p - 1) \left[2(1 - \mu)^2 - 2(1 - \mu)^{K-2} \right] \geq 0. \end{aligned}$$

Case 2: Suppose that there is a “pool” of size $K = 1$ (i.e., suppose that one of the subjects is individually tested). Because N is even, this implies that there must be at least one other pool with odd size. From the previous case analyzed, it will never be optimal to have a pool of size $K \geq 4$. So, we only need to consider two possibilities.

Case 2.1: Suppose that there is another “pool” of size $K = 1$ as well. In this case, the expected number of false positives from these two individual tests is given by

$$2(1 - \mu)(1 - S_p). \quad (18)$$

Meanwhile, if we were to pool these two specimens together to form a pool of size 2, they would generate the following expected number of false positives

$$2 \left[(1 - \mu)(1 - S_p)S_e - (1 - \mu)^2(1 - S_p)(S_e + S_p - 1) \right]. \quad (19)$$

Subtracting (18) by (19) yields

$$2(1 - \mu)^2(1 - S_p)(S_e + S_p - 1) + 2(1 - \mu)(1 - S_p)(1 - S_e) > 0.$$

Case 2.2: Suppose that the other pool is of size 3. Then, the expected number of false positives from these two pools is given by

$$(1 - \mu)(1 - S_p) + 3 \left[(1 - \mu)(1 - S_p)S_e - (1 - \mu)^3(1 - S_p)(S_e + S_p - 1) \right]. \quad (20)$$

If, however, we randomly took one of the subjects from the pool of size 3 and matched him with the subject being individually tested, these two pools would generate the following expected number of false positives

$$4 \left[(1 - \mu)(1 - S_p)S_e - (1 - \mu)^2(1 - S_p)(S_e + S_p - 1) \right]. \quad (21)$$

Subtracting (20) by (21) yields:

$$(1 - \mu)(1 - S_p)(1 - S_e) + (1 - \mu)^2 [4 - 3(1 - \mu)] (1 - S_p)(S_e + S_p - 1) \geq 0.$$

Case 3: The only case remaining occurs when one of the pools has size $K = 3$. Suppose by way of contradiction that this pool configuration is optimal. From the previous cases, optimality requires that at least one other pool also has size $K = 3$. Together, these two pools would generate the following expected number of false positives:

$$6 [(1 - \mu)(1 - S_p)S_e - (1 - \mu)^3(1 - S_p)(S_e + S_p - 1)]. \quad (22)$$

If, however, we randomly matched the subjects from these pools to create 3 pools of size 2, the expected number of false positives from these pools would be given by

$$6 [(1 - \mu)(1 - S_p)S_e - (1 - \mu)^2(1 - S_p)(S_e + S_p - 1)]. \quad (23)$$

Subtracting (22) by (23) yields

$$(1 - S_p)(S_e + S_p - 1)(1 - \mu)^2 \mu \geq 0.$$

Ordered Pooling: Suppose that ordered pooling is implemented.

Case 1: Aprahamian, Bish and Bish [2019] has already shown that, when our objective is to minimize the expected number of false positives, it is never optimal to have a pool $K \geq 4$.

Case 2: Suppose by way of contradiction that we have a pool configuration that minimizes the expected number of false positives, where one of the “pools” has size $K = 1$ (i.e., one of the subjects is individually tested). Let q_i be the probability of infection of this subject that is individually tested.

From Case 1 and the fact that N is even, there must be at least one other pool of size $K = 3$ or size $K = 1$.

Case 2.1: Suppose that there is another “pool” of size $K = 1$, and suppose that the person in this pool has probability of infection q_j . Then, the expected number of false positives from these two pools is given by

$$(1 - q_i)(1 - S_p) + (1 - q_j)(1 - S_p).$$

Meanwhile, if these two subjects were grouped together to form a pool of size $K = 2$, they would generate the following expected number of false positives

$$\begin{aligned}
& (1 - q_i)(1 - S_p) \underbrace{[(1 - q_j)(1 - S_p) + q_j S_e]_{\leq S_e}} \\
& + (1 - q_j)(1 - S_p) \underbrace{[(1 - q_i)(1 - S_p) + q_i S_e]_{\leq S_e}} \leq \\
& (1 - q_i)(1 - S_p) S_e + (1 - q_j)(1 - S_p) S_e \leq \\
& (1 - q_i)(1 - S_p) + (1 - q_j)(1 - S_p).
\end{aligned}$$

Case 2.2: Suppose by way of contradiction that there is a pool of size $K = 3$.

Without loss of generality, suppose that the probabilities of infection from subjects in this pool is given by (q_1, q_2, q_3) . Then, the expected number of false positives from this pool plus the probability of a false positive from the subject that is individually tested is given by

$$\begin{aligned}
& (1 - q_i)(1 - S_p) + \sum_{j=1}^3 (1 - q_j)(1 - S_p) \left[\prod_{l \in \{1,2,3\} \setminus j} (1 - q_l)(1 - S_p) \right. \\
& \left. + S_e \left(1 - \prod_{l \in \{1,2,3\} \setminus j} (1 - q_l) \right) \right].
\end{aligned}$$

Meanwhile, if we were to merge together these two pools to generate a pool of size $K = 2$, these subjects would generate the following expected number of

false positives:

$$\begin{aligned}
& \sum_{j \in \{1,2,3,i\}} (1 - q_j)(1 - S_p) \left[\prod_{l \in \{1,2,3,i\} \setminus j} (1 - q_l)(1 - S_p) \right. \\
& \quad \left. + S_e \left(1 - \prod_{l \in \{1,2,3,i\} \setminus j} (1 - q_l) \right) \right] \leq \\
& (1 - q_i)(1 - S_p)S_e + \sum_{j=1}^3 (1 - q_j)(1 - S_p) \left[\prod_{l \in \{1,2,3\} \setminus j} (1 - q_l)(1 - S_p) \right. \\
& \quad \left. + S_e \left(1 - \prod_{l \in \{1,2,3\} \setminus j} (1 - q_l) \right) \right] \leq \\
& (1 - q_i)(1 - S_p) + \sum_{j=1}^3 (1 - q_j)(1 - S_p) \left[\prod_{l \in \{1,2,3\} \setminus j} (1 - q_l)(1 - S_p) \right. \\
& \quad \left. + S_e \left(1 - \prod_{l \in \{1,2,3\} \setminus j} (1 - q_l) \right) \right].
\end{aligned}$$

Case 3: The only case remaining to analyze occurs when we have two pools of size $K = 3$.

Without loss of generality, suppose that the probabilities of infection from the first of these pools is given by (q_1, q_2, q_3) , and that the probabilities of infection from the other pool of size $K = 3$ is given by (q_4, q_5, q_6) . Then, the expected number of false positives from these two pools is given by

$$\begin{aligned}
& \sum_{j=1}^3 (1 - q_j)(1 - S_p) \left[\prod_{l \in \{1,2,3\} \setminus j} (1 - q_l)(1 - S_p) + S_e \left(1 - \prod_{l \in \{1,2,3\} \setminus j} (1 - q_l) \right) \right] \\
& + \sum_{j=4}^6 (1 - q_j)(1 - S_p) \left[\prod_{l \in \{4,5,6\} \setminus j} (1 - q_l)(1 - S_p) + S_e \left(1 - \prod_{l \in \{4,5,6\} \setminus j} (1 - q_l) \right) \right].
\end{aligned}$$

If, however, we were to rearrange these two pools to form 3 pools of size $K = 2$, where the subjects with probabilities of infection q_1 and q_2 were matched together; those with probabilities of infection q_3 and q_4 were matched together; and those with probabilities of infection q_5 and q_6 were matched together, we would end up with the following expected number of false positives:

$$\begin{aligned}
& (1 - q_1)(1 - S_p) [(1 - q_2)(1 - S_p) + q_2 S_e] \\
& + (1 - q_2)(1 - S_p) [(1 - q_1)(1 - S_p) + q_1 S_e] \\
& + (1 - q_3)(1 - S_p) [(1 - q_4)(1 - S_p) + q_4 S_e] \\
& + (1 - q_4)(1 - S_p) [(1 - q_3)(1 - S_p) + q_3 S_e] \\
& + (1 - q_5)(1 - S_p) [(1 - q_6)(1 - S_p) + q_6 S_e] \\
& + (1 - q_6)(1 - S_p) [(1 - q_5)(1 - S_p) + q_5 S_e] \leq \\
& (1 - q_1)(1 - S_p) [(1 - q_2)(1 - q_3)(1 - S_p) + S_e[1 - (1 - q_2)(1 - q_3)]] \\
& + (1 - q_2)(1 - S_p) [(1 - q_1)(1 - q_3)(1 - S_p) + S_e[1 - (1 - q_1)(1 - q_3)]] \\
& + (1 - q_3)(1 - S_p) [(1 - q_1)(1 - q_2)(1 - S_p) + S_e[1 - (1 - q_1)(1 - q_2)]] \\
& + (1 - q_4)(1 - S_p) [(1 - q_5)(1 - q_6)(1 - S_p) + S_e[1 - (1 - q_5)(1 - q_6)]] \\
& + (1 - q_5)(1 - S_p) [(1 - q_4)(1 - q_6)(1 - S_p) + S_e[1 - (1 - q_4)(1 - q_6)]] \\
& + (1 - q_6)(1 - S_p) [(1 - q_4)(1 - q_5)(1 - S_p) + S_e[1 - (1 - q_4)(1 - q_5)]] = \\
& \sum_{j=1}^3 (1 - q_j)(1 - S_p) \left[\prod_{l \in \{1,2,3\} \setminus j} (1 - q_l)(1 - S_p) + S_e \left(1 - \prod_{l \in \{1,2,3\} \setminus j} (1 - q_l) \right) \right] \\
& + \sum_{j=4}^6 (1 - q_j)(1 - S_p) \left[\prod_{l \in \{4,5,6\} \setminus j} (1 - q_l)(1 - S_p) + S_e \left(1 - \prod_{l \in \{4,5,6\} \setminus j} (1 - q_l) \right) \right].
\end{aligned}$$

■

A.3 Expected Number of False Negatives

For each group G_g and each $j \in G_g$ we define

$$\begin{aligned}
A_{G_g, j} &\equiv \sum_{I=0}^{K-1} (I+1)(1 - h(I+1)S_e) P_{G_g \setminus \{j\}}(I), \\
B_{G_g, j} &\equiv \sum_{I=0}^{K-1} I(1 - h(I)S_e) P_{G_g \setminus \{j\}}(I).
\end{aligned}$$

From the above expressions we have that $A_{G_g, j}$ corresponds to the expected number of false negatives in group G_g conditional that $j \in G_g$ is infected. Similarly, $B_{G_g, j}$ is the expected number of false negatives in group G_g conditional that $j \in G_g$ is not infected.

Now notice that, for any $j \in G_g$,

$$fn_g = q_j A_{G_g, j} + (1 - q_j) B_{G_g, j} = q_j (A_{G_g, j} - B_{G_g, j}) + B_{G_g, j}.$$

Lemma A.7 (*Online Appendix from Saraiva [2023]*) *If hypothesis 3 holds we have that, for each group G_g and each $l \in G_g$,*

$$A_{G_g, l} - B_{G_g, l}$$

is decreasing in the probability of infection from each subject in $G_g \setminus \{l\}$.

Proof of proposition 8: Analogous to the proof of proposition 5 but, instead of invoking Lemma A.5, we invoke Lemma A.7. ■

Proof of theorem 3: Proof omitted, as it is analogous to the proof of theorems 1 and 2. ■

Lemma A.8 *For each $\mu \in [a, b]$ denote G_μ as a group with $K - 1$ subjects, each of which has probability of infection μ . Also, let*

$$A_{G_\mu} \equiv \sum_{I=0}^{K-1} (I+1)(1-h(I+1)S_e) \binom{K-1}{I} \mu^I (1-\mu)^{K-1-I}$$

$$B_{G_\mu} \equiv \sum_{I=0}^{K-1} I(1-h(I)S_e) \binom{K-1}{I} \mu^I (1-\mu)^{K-1-I}$$

i.e., A_{G_μ} is the expected number of false negatives from group G_μ conditional that one of its subjects is infected, whereas B_{G_μ} is the expected number of false negatives from group G_μ conditional that one of its subjects is not infected.

Then, we have that

$$N \sum_{I=0}^{K-1} (1 - S_e h(I+1)) \binom{K-1}{I} [\mu^{I+1} (1-\mu)^{K-1-I}] = n [\mu A_{G_\mu} + (1-\mu) B_{G_\mu}].$$

Proof: Proof omitted, as it is similar to the proof of Lemma A.6. ■

Proof of proposition 9: Analogous to the proof of proposition 6, but, instead of invoking Lemmas A.5 and A.6, we invoke Lemmas A.7 and A.8. ■

A.4 Upper bounds with partial information on μ , a , b , S_e and S_p

Proof of corollary 7: To compute $UB_o^T(\mu, n, K)$, $UB_o^{FP}(\mu, n, K)$ and $UB_o^{FN}(\mu, n, K)$ we basically choose a probability of infection configuration q that maximizes an objective function subject to the constraint that $\sum_i q_i = \mu$ and $q_i \in [a, b]$ for all i . When we decrease a or increase b , we expand the domain of our objective function, without altering

our objective function, which implies that a decrease in a or an increase in b should never reduce our upper bounds. ■

Proof of proposition 11: Because the function $f(x) \equiv x^K$ is convex and because $(1 - \mu) = \alpha(1 - a) + (1 - \alpha)(1 - b)$, we have that the term $[\alpha(1 - a)^K + (1 - \alpha)(1 - b)^K - (1 - \mu)^K]$ from equations (6) and (10) is greater than or equal to zero. By analyzing the first derivatives of (6) and (10) with respect to S_e and S_p we can see that expression (6) is strictly increasing in S_e and S_p , while expression (10) is strictly increasing in S_e and strictly decreasing in S_p when $2(1 - S_p) < S_e$.

One can clearly see that the term

$$[\alpha(1 - a)^K + (1 - \alpha)(1 - b)^K - (1 - \mu)^K] = \left[\frac{b - \mu}{b - a}(1 - a)^K + \frac{\mu - a}{b - a}(1 - b)^K - (1 - \mu)^K \right]$$

is strictly concave with respect to μ . Therefore, from the first order condition we have that the unique maximum is achieved at

$$\mu^* = 1 - \left[\frac{(1 - a)^K - (1 - b)^K}{K(b - a)} \right]^{1/(K-1)}.$$

■

B How large N needs to be for the upper bounds to be accurate?

To derive the upper bounds presented in the paper, we had to assume the batch size, N , to be sufficiently large. This is done mainly to ensure that the average prevalence from the batch $\bar{q}$ can be approximated by the population prevalence, μ . But precisely how large N has to be for the approximation to be accurate?

In this section we show how one can find the minimum level of N that ensures that $\bar{q}$ always stays sufficiently close to μ with a high probability. To accomplish this goal, we 1) first find the maximum level of variance that $\bar{q}$ can possibly achieve, and set a tolerance level $\tau > 0$ for the deviations that $\bar{q}$ can have over μ . 2) Using this maximum variance, we find the maximum batch size N such that $Prob(|\bar{q} - \mu| \leq \tau) \geq 0.95$. 3) Then, we find

$$\max_{\tilde{\mu} \in [\mu - \tau, \mu + \tau]} \lim_{n \rightarrow \infty} UB_o^T(\tilde{\mu}, n, K)$$

and check how far away this optimum is from $\lim_{n \rightarrow \infty} UB_o^T(\mu, n, K)$. If the difference is very small, this means that the N that we found in the first step is indeed large enough so that we can be confident that our upper bound is reliable.

It should be noticed that if, after implementing this process, the minimum N required to ensure that our upper bound is accurate is unrealistically large, one can still get a conservative upper bound by maximizing the asymptotic upper bound with respect to μ on the entire domain $[a, b]$, as discussed in section 7. It should also be noticed that the minimum N that we obtain by implementing this process tends to be overly conservative, as it relies on a distribution that generates the highest possible noise in $\bar{q}$. Finally, it should also be noticed that if a laboratory intends to test multiple batches, deviations of $\bar{q}$ from μ will tend to average themselves out, so that this analysis becomes less relevant.

The next lemma shows how to implement step 1) from the process we just described.

Lemma B.1 *Let $Prob(p_i)$ be the probability that a subject in the population has probability of infection p_i , so that $\sum_i Prob(p_i) = 1$. Then, the configuration of probabilities of infection that maximize the variance of $\bar{q}$ subject to the constraint that $\mathbb{E}(\bar{q}) = \mu$ and that $p_i \in [a, b]$ for every p_i such that $Prob(p_i) > 0$ is given by $Prob(a) = \alpha = \frac{b-\mu}{b-a}$ and $Prob(b) = 1 - \alpha = \frac{\mu-a}{b-a}$, which yields the following variance for $\bar{q}$:*

$$\sigma^2 = \alpha \frac{(a - \mu)^2}{n} + (1 - \alpha) \frac{(b - \mu)^2}{n}. \quad (24)$$

Proof: Suppose by way of contradiction that there exists a p_i and p_j such that $a < p_i \leq p_j < b$ and $Prob(p_i) > 0$ and $Prob(p_j) > 0$. Then, there is a ε sufficiently small such that $p_i - \frac{\varepsilon}{Prob(p_i)} > a$ and $p_j + \frac{\varepsilon}{Prob(p_j)} < b$. Then, replacing p_i and p_j by $p_i - \frac{\varepsilon}{Prob(p_i)} > a$ and $p_j + \frac{\varepsilon}{Prob(p_j)} < b$, respectively, would preserve the mean of the distribution, μ , while not violating the boundary constraint that $p_i, p_j \in [a, b]$, and this would clearly increase the variance of $\bar{q}$. ■

Lemma B.1 in combination with the central limit theorem (CLT) implies that $\bar{q}$ can be approximated by a normal distribution with mean μ and an unknown variance that is less than or equal to $\alpha \frac{(a-\mu)^2}{n} + (1 - \alpha) \frac{(b-\mu)^2}{n}$. Therefore, for a given tolerance level $\tau > 0$,

we have that¹

$$Prob(|\bar{q} - \mu| \leq \tau) \geq 2\Phi \left(\tau \frac{\sqrt{N}}{\sqrt{\alpha(a - \mu)^2 + (1 - \alpha)(b - \mu)^2}} \right) - 1, \quad (25)$$

where Φ corresponds to the cdf from a standard normal distribution.

Using simple numerical tools one can find the value of N such that

$$2\Phi \left(\tau \frac{\sqrt{N}}{\sqrt{\alpha(a - \mu)^2 + (1 - \alpha)(b - \mu)^2}} \right) - 1 = 0.95.$$

We then optimize

$$\max_{\tilde{\mu} \in [\mu - \tau, \mu + \tau]} \lim_{n \rightarrow \infty} UB_o^T(\tilde{\mu}, n, K)$$

and check how far away this maximum is from $\lim_{n \rightarrow \infty} UB_o^T(\mu, n, K)$.

Example B.1 Suppose that $\mu = 0.01$, $a = 0.001$, $b = 0.05$, $K = 4$ and that the dilution function is given by equation (1) with $S_e = S_p = 0.99$ and $\delta = 0.04$. Then, if we apply a tolerance level of $\tau = .005$, we have that, as long as $N > 55$,

$$Prob(|\bar{q} - \mu| \leq \tau) \geq 0.95.$$

Moreover, if we assume that the cost of a test is given by 30 USD, we have that the maximum savings associated with the expected number of tests at $\mu = 0.01$ would be 0.056 USD, i.e., approximately 5 and a half cents per subject (the reader is invited to replicate this result using our Shiny app https://rt0c5u-gustavo-saraiva.shinyapps.io/shiny_app/). Meanwhile,

$$\max_{\tilde{\mu} \in [\mu - \tau, \mu + \tau]} \lim_{n \rightarrow \infty} UB_o^T(\tilde{\mu}, n, K) = 0.076,$$

so that, with a 95% confidence level, our upper bound to the savings in the expected number of tests associated with ordered pooling can be “far off from the truth” by at most 2 cents. In this case, if we wish to be more conservative regarding our estimation of the maximum benefits of implementing ordered pooling, we should assume 0.076 USD to be our upper bound, instead of 0.056 USD.

This result is illustrated in figure I. The blue dot from the figure corresponds to our upper bound (equation (5)) evaluated at $\mu = 0.01$. The gray area corresponds to the levels that $\bar{q}$ and our upper bound assume at least 95% of the time, assuming that $N > 55$.

¹A looser condition that does not rely on the CLT can be obtained by using Chebyshev’s inequality,

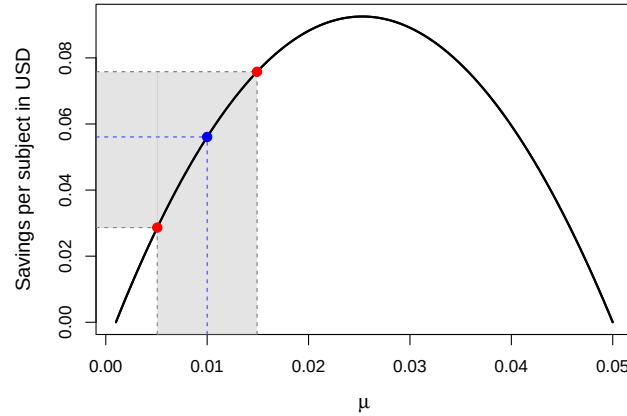


Figure I: The savings associated with ordered pooling as a function of μ , assuming that $a = 0.001$, $b = 0.05$, $K = 4$ and that the dilution function is given by equation (1) with $S_e = S_p = 0.99$ and $\delta = 0.04$. The blue dot corresponds to the benefits of implementing ordered pooling evaluated at the population prevalence $\mu = 0.01$. Assuming a tolerance level for $\bar{q}$ of $\tau = 0.005$, we have that our upper bound can be off by at most 2 cents

C Alternative Parametrization for the Chlamydia Case Study

To be consistent with previous research and because of some limitations from the dataset available, some arbitrary assumptions were made in the paper regarding some of the parameters of the model for the Chlamydia Case Study. In particular, it was assumed that the dilution parameter δ from $h(I, K)$ was equal to 0.15, in spite of some evidence suggesting that this parameters should be much lower.

From figure II, however, we can see that the savings in costs associated with ordered pooling do not change very much if we assume lower values for δ . If anything, our upper bounds slightly diminish as we decrease δ , as decreasing this parameter decreases the dilution effect: in the limit, when δ is very low, both ordered pooling and random pooling generate the same expected costs associated with false negative results (see proposition 10).

which implies that

$$Prob(|\bar{q} - \mu| \leq \tau) \geq \frac{\alpha(a - \mu)^2 + (1 - \alpha)(b - \mu)^2}{N\tau^2}.$$

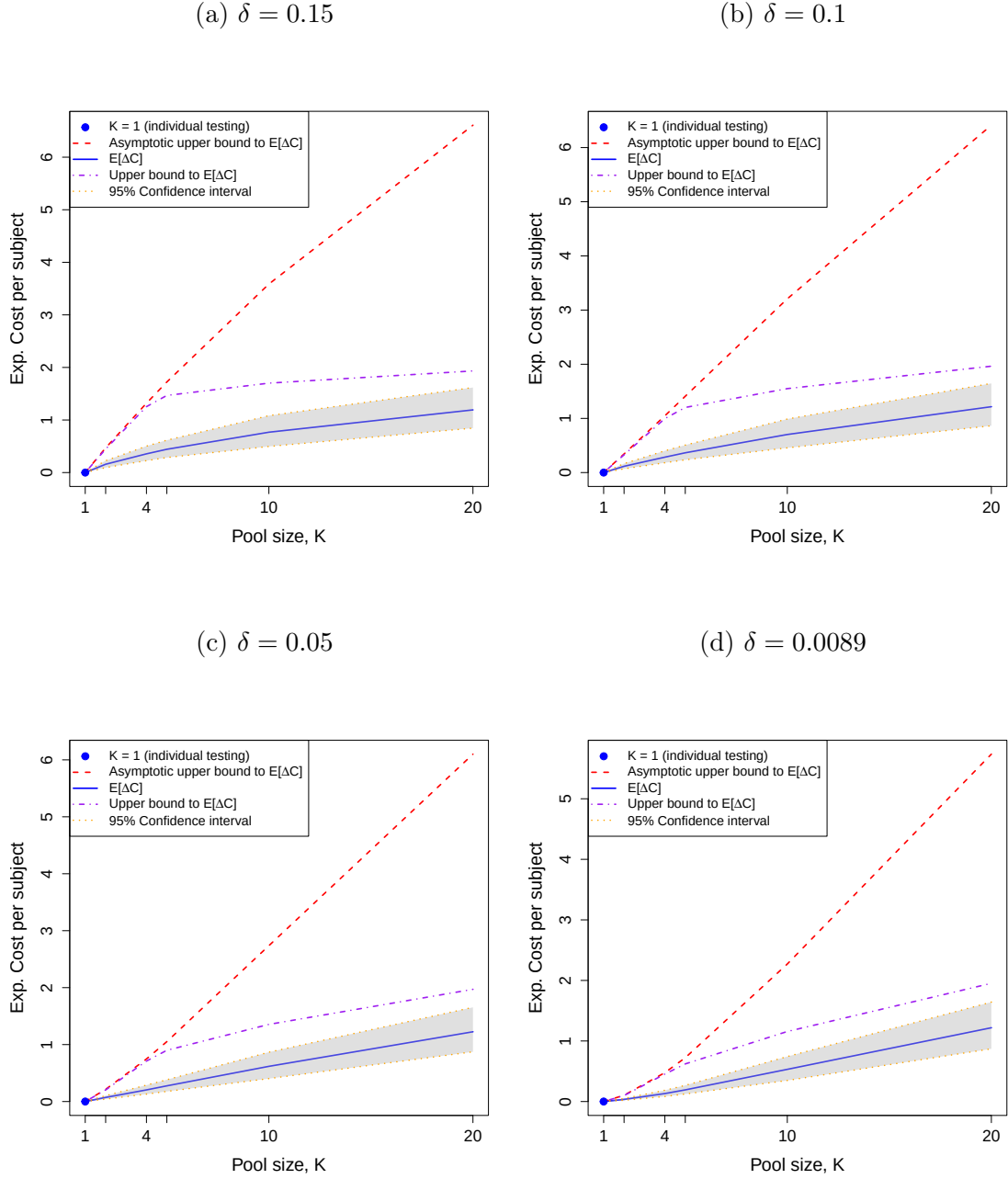


Figure II: Expected benefits per subject of implementing ordered pooling for Chlamydia, assuming $N = 100$. The gray area corresponds to a 95% confidence interval using non-parametric bootstrap and assuming that we test 10 batches in total. The blue dot corresponds to the case in which subjects are individually tested, i.e. $K = 1$. Parameters for the dilution function (16): $S_e = .99$, $S_p = .98$ and $\delta \in \{.15, .1, .05, .0089\}$. Parameters for the upper bounds: $\mu = 0.0097$, $a = 0.0017$ and $b = 0.1919$.

Another arbitrary assumption made corresponds to the underreporting factor of 3 used. As explained in the text, assuming a higher underreporting factor tends to substantially increase the expected benefits of implementing ordered pooling. Figure III depicts the benefits of reporting ordered pooling when assuming an underreporting factor of 3 vs. an underreporting factor of 1.

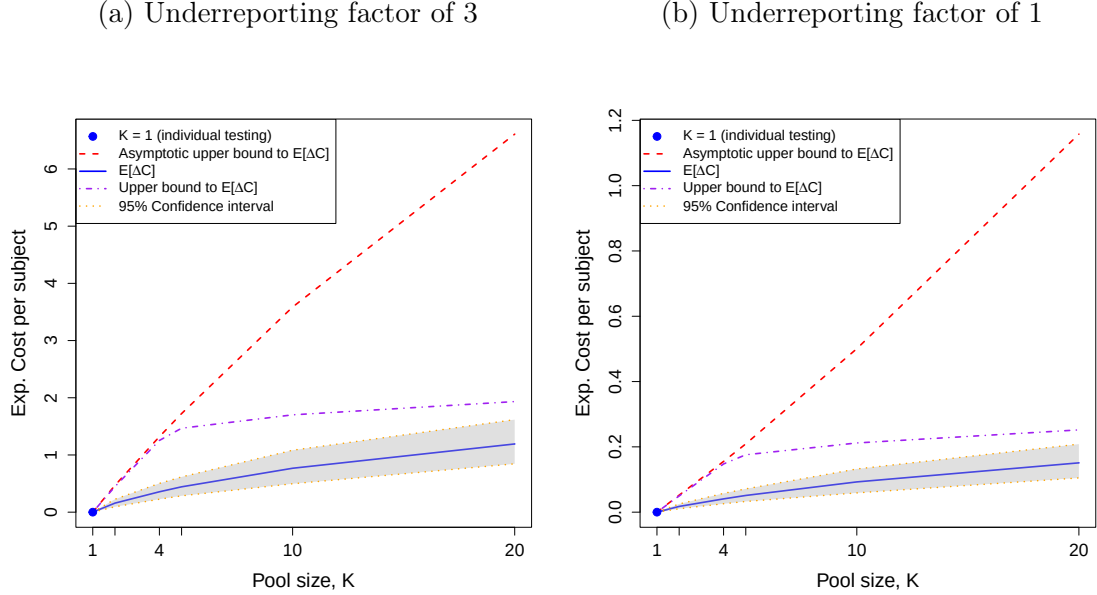


Figure III: Expected benefits per subject of implementing ordered pooling for Chlamydia, assuming $N = 100$. The gray area corresponds to a 95% confidence interval using non-parametric bootstrap and assuming that we test 10 batches in total. The blue dot corresponds to the case in which subjects are individually tested, i.e. $K = 1$. Parameters for the dilution function (16): $S_e = .99$, $S_p = .98$ and $\delta = .15$.

Notice that these results can be easily replicated with the assistance of the companion app https://rt0c5u-gustavo-saraiva.shinyapps.io/shiny_app/ developed for this paper.

D Upper bounds without information on the dilution effect

In this section, we derive upper bounds to the benefits of implementing ordered pooling without relying on information regarding the dilution function h . The trade-off is that

these upper bounds also tend to be much looser than the ones derived in the main text. But as we will see at the end of this section, these upper bounds may still be useful in cases in which the prevalence is very small, as, in these cases, even these loose upper bounds predict small gains in implementing ordered pooling as opposed to random pooling.

Though we do not need to have a specific formula for the dilution effect to derive the upper bounds from this section, small constraints must be imposed on the dilution effect. More specifically, throughout this section we will assume that $h(K, K) = S_e$, i.e., the probability of detecting an infection in a pooled sample comprised entirely of infected specimens is the same as the probability of detecting an infection from an individual sample of an infected subject. Similarly, it is also natural to assume that $h(0, K) = 1 - S_p$. As an example, the dilution function

$$h(I, K) = (1 - S_p) + (S_p + S_e - 1) \left(\frac{I}{K} \right)^\delta,$$

with $\delta \geq 0$ used in Aprahamian, Bish and Bish [2018] and Saraiva [2023] satisfy these properties. Similarly, the dilution function

$$h(I, K) = \begin{cases} 1 - S_p, & \text{if } I = 0 \\ S_e, & \text{else} \end{cases},$$

which corresponds to the case in which pooled testing is not subject to any dilution effect whatsoever, also satisfy these properties.

Assumption D.1 $h(0, K) = 1 - S_p$ and $h(K, K) = S_e$.²

D.1 Upper bound to the Expected Number of Tests without information regarding the dilution effect

In this section we show how to compute an upper bound to the reduction in the expected number of tests obtained when implementing ordered pooling as opposed to random pooling without relying on information regarding the dilution effect. For a given pool size K , we show that only four parameters are required to compute the upper bound: μ , a , S_e and S_p .

In this and in the following section we will assume that $\mu \leq \frac{1}{K}$, i.e., the prevalence of the disease is not too high. In many practical applications this is a very mild assumptions for the following reasons:

²This assumption can be relaxed to $h(0, K) \geq 1 - S_p$ and $h(K, K) \leq S_e$.

1. Pooled testing is only cost-effective when the prevalence of the disease is small (e.g., Dorfman [1943]). Indeed, when the prevalence of the disease is high, most pools end up infected, and therefore, retested, in which case laboratories are better off just testing everyone individually.
2. In many practical applications pool sizes are rather small to mitigate the negative impact that dilution effects may have on the precision of the test. As an example, throughout the COVID-19 pandemic, Chile relied heavily on pooled testing, and in the great majority of cases, pool sizes were no greater than $K = 5$ (Basso et al. [2023]). Meanwhile, in Kenya samples for RT-PCR tests were usually pooled into groups of size 6 (Agoti et al. [2021]).
3. The optimum pool size tends to decrease with the prevalence level, thus guaranteeing that $\mu \leq \frac{1}{K}$ (see, for example, table I from Dorfman [1943] or the shiny application from Christopher Bilder available at his personal website <https://www.chrisbilder.com/shiny>).

Assumption D.2 $\mu \leq \frac{1}{K}$.

For our datasets and the pool sizes we consider, assumption D.2 always holds. Moreover, this assumption can be relaxed by the milder assumption that $\mu(1 - \mu)^{K-1} \geq a(1 - a)^{K-1}$.

Proposition D.3 *Upper bound to savings in the expected number of tests: If assumptions 1, D.2 and D.1 hold, we have that*

$$\frac{T_r^K(\bar{q}, n, K) - T_o^K(q_1, q_2, \dots, q_N)}{N} \leq (S_p + S_e - 1) [(1 - a)^K - (1 - \mu)^K].$$

Proof: Because the probability that each subject is retested in a group is increasing in the subject's probability of infection and also on the probability of infection of everyone else within the group,³ we clearly have that a lower bound to the expected number of test when implementing an ordered partition in which all pools have size K is given by computing the same expression as before, but assuming that all subjects have the lowest

³This happens because $h(I, K)$ is increasing in I , and because increasing the probability of infection of someone within a pool causes a shift in the distribution of I that first order-stochastically dominates the previous distribution.

possible probability of infection a . This yields the following expression:

$$\underline{T}_o^K \equiv \frac{N}{K} + \frac{N}{K} K \left[\sum_{I=0}^K h(I, K) \binom{K}{I} a^I (1-a)^{K-I} \right]. \quad (26)$$

So an upper bound to the expected reduction in the number of tests per subject that one can achieve by implementing ordered pooling as opposed to random pooling is given by:

$$\frac{T_r^K(\bar{q}, n, K) - \underline{T}_o^K}{N} = \sum_{I=0}^K h(I, K) \binom{K}{I} [\mu^I (1-\mu)^{K-I} - a^I (1-a)^{K-I}]. \quad (27)$$

Because $h(0, K) = 1 - S_p$, we can rewrite equation (27) as

$$\begin{aligned} \frac{T_r^K(\bar{q}, n, K) - \underline{T}_o^K}{N} &= (1 - S_p)((1-\mu)^K - (1-a)^K) \\ &\quad + \sum_{I=1}^K h(I, K) \binom{K}{I} [\mu^I (1-\mu)^{K-I} - a^I (1-a)^{K-I}]. \end{aligned} \quad (28)$$

Because $\mu \geq a$, we have that $\mu^I (1-\mu)^{K-I} - a^I (1-a)^{K-I} \geq 0$ for all $I \geq 1$ if and only if $\mu(1-\mu)^{K-1} - a(1-a)^{K-1} \geq 0$. Noticing that the function $f(q) \equiv q(1-q)^{K-1}$ is unimodal within the interval $[0, 1]$, having a single peak at $q^* = \frac{1}{K}$, and noticing that $a < \mu$, we have that $\mu \leq \frac{1}{K}$ implies that $\mu(1-\mu)^{K-1} - a(1-a)^{K-1} \geq 0$, which in turn implies that $\mu^I (1-\mu)^{K-I} - a^I (1-a)^{K-I} \geq 0$ for all $I \geq 1$. So, for every $I \geq 1$, the term $h(I, K)$ from expression (28) is always being multiplied by a non-negative term, which implies that this expression is increasing in $h(I, K)$ for all $I \geq 1$. But from assumption D.1 we have that $h(I, K) \leq S_e$. Therefore

$$\begin{aligned} \frac{T_r^K(\bar{q}, n, K) - \underline{T}_o^K}{N} &\leq (1 - S_p)((1-\mu)^K - (1-a)^K) \\ &\quad + S_e \left[\underbrace{\sum_{I=1}^K \binom{K}{I} \mu^I (1-\mu)^{K-I}}_{=1-(1-\mu)^K} - \underbrace{\sum_{I=1}^K \binom{K}{I} a^I (1-a)^{K-I}}_{=1-(1-a)^K} \right] \\ &\leq (S_p + S_e - 1) [(1-a)^K - (1-\mu)^K]. \end{aligned}$$

■

Notice that, if the tester does not know a , only the overall prevalence μ , he can conservatively set $a = 0$ to get a looser upper bound to the benefits of implementing ordered pooling in terms of minimizing the expected number of tests:

$$\frac{T_r^K(\bar{q}, n, K) - T_o^K(q_1, \dots, q_N)}{N} \leq (S_p + S_e - 1) [1 - (1 - \bar{q})^K].$$

Also notice that, the smaller μ is, the smaller is our upper bound to $(T_r^K - T_o^K)/N$. Intuitively, this happens because, when the prevalence is small, very few subjects are retested, regardless of how they are matched to form the pools, in which case, one should expect small gains of implementing ordered pooling as opposed to random pooling. The actual benefits of implementing ordered pooling are not, however, monotonic in μ . Indeed, in the extreme cases in which $\mu = 0$ or $\mu = 1$, all subjects have the same probability of infection, in which case ordered pooling becomes equivalent to random pooling. In practice however, the prevalence of the disease tends to be relatively small, and therefore in the region in which the benefits of implementing ordered pooling is increasing in μ .

D.2 Upper bound to the Expected Number of False Positives without information regarding the dilution effect

In this section we show how to compute an upper bound to the reduction in the expected number of false positives obtained when implementing ordered pooling as opposed to random pooling, without relying on information regarding the dilution effect. Like in the previous section, computing the upper bound only requires knowing the pool size K and estimates for the following parameters: μ , a , S_e and S_p .

Approximating $\bar{q}$ by μ (a good approximation when N is large) allows us to derive a simple formula to an upper bound to $\frac{FP_r^K(\mu, n, K) - FP_o^K(q_1, q_2, \dots, q_N)}{N}$.

Proposition D.4 *Upper bound to the reduction in the expected number of false positives: If assumptions 1, D.2 and D.1 hold, then, for a given realization of $(q_1, q_2, \dots, q_N)$ with $\sum_i q_i/N = \mu$ we have that*

$$\frac{FP_r^K(\mu, n, K) - FP_o^K(q_1, \dots, q_N)}{N} \leq (1 - S_p)(1 - \mu)(S_p + S_e - 1) [(1 - a)^{K-1} - (1 - \mu)^{K-1}].$$

Proof: Clearly, the probability that one receives a false positive classification is increasing in the probability of infection from everyone else within the pool. So the following expression corresponds to a lower bound to the expected number of false positives that

one gets under ordered pooling:

$$\begin{aligned}
\underline{FP}_o^K &= \sum_{i=1}^N (1 - q_i) \left[\sum_{I=0}^{K-1} h(I, K) \binom{K-1}{I} q_1^I (1 - q_1)^{K-1-I} \right] (1 - S_p) \\
&= N(1 - \bar{q}) \left[\sum_{I=0}^{K-1} h(I, K) \binom{K-1}{I} q_1^I (1 - q_1)^{K-1-I} \right] (1 - S_p) \\
&= N(1 - \mu) \left[\sum_{I=0}^{K-1} h(I, K) \binom{K-1}{I} \mu^I (1 - \mu)^{K-1-I} \right] (1 - S_p)
\end{aligned}$$

We wish to find an upper bound to

$$\begin{aligned}
\frac{FP_r^K - \underline{FP}_o^K}{N} &= (1 - S_p)(1 - \mu) \left[\sum_{I=0}^{K-1} h(I, K) \binom{K-1}{I} (\mu^I (1 - \mu)^{K-1-I} \right. \\
&\quad \left. - a^I (1 - a)^{K-1-I}) \right]. \tag{29}
\end{aligned}$$

From assumption D.1 we have that $h(0, K) = 1 - S_p$, so we can rewrite equation (29) as

$$\begin{aligned}
\frac{FP_r^K - \underline{FP}_o^K}{N} &= (1 - S_p)(1 - \mu) \left[(1 - S_p) ((1 - \mu)^{K-1} - (1 - a)^{K-1}) + \right. \\
&\quad \left. + \sum_{I=1}^{K-1} h(I, K) \binom{K-1}{I} (\mu^I (1 - \mu)^{K-1-I} - a^I (1 - a)^{K-1-I}) \right]. \tag{30}
\end{aligned}$$

Because $\mu \geq a$, we have that $\mu^I (1 - \mu)^{K-1-I} - a^I (1 - a)^{K-1-I} \geq 0$ for all $I \geq 1$ if and only if $\mu(1 - \mu)^{K-2} - a(1 - a)^{K-2} \geq 0$. Noticing that the function $f(q) \equiv q(1 - q)^{K-2}$ is unimodal within the interval $[0, 1]$, having a single peak at $q^* = \frac{1}{K-1}$, and noticing that $a \leq \mu$, we have that $\mu \leq \frac{1}{K} < q^* = \frac{1}{K-1}$ implies that $\mu(1 - \mu)^{K-2} - a(1 - a)^{K-2} \geq 0$, which in turn implies that $\mu^I (1 - \mu)^{K-1-I} - a^I (1 - a)^{K-1-I} \geq 0$ for all $I \geq 1$. So, for every $I \geq 1$, the term $h(I, K)$ from expression (30) is always being multiplied by a non-negative term, which implies that this expression is increasing in $h(I, K)$ for all $I \geq 1$. But from assumption D.1 we have that $h(I, K) \leq S_e$. Therefore,

$$\begin{aligned}
\frac{FP_r^K - \underline{FP}_o^K}{N} &\leq (1 - S_p)(1 - \mu) \left[(1 - S_p) ((1 - \mu)^{K-1} - (1 - a)^{K-1}) + \right. \\
&\quad \left. + S_e \sum_{I=1}^{K-1} \binom{K-1}{I} (\mu^I (1 - \mu)^{K-1-I} - a^I (1 - a)^{K-1-I}) \right] \\
&\leq (1 - S_p)(1 - \mu) [(1 - S_p) ((1 - \mu)^{K-1} - (1 - a)^{K-1}) \\
&\quad + S_e (1 - (1 - \mu)^{K-1} - (1 - (1 - a)^{K-1}))] \\
&\leq (1 - S_p)(1 - \mu)(S_p + S_e - 1) [(1 - a)^{K-1} - (1 - \mu)^{K-1}].
\end{aligned}$$

■

Like in the previous section, if the tester does not know a , only the overall prevalence μ , he can conservatively set $a = 0$ to get a looser upper bound to the benefits of implementing ordered pooling in terms of reducing the expected number of false positives:

$$\frac{FP_r^K - FP_o^K}{N} \leq (1 - S_p)(1 - \mu)(S_p + S_e - 1) [1 - (1 - \mu)^{K-1}].$$

D.3 Upper bound to the Expected Number of False Negatives with partial information regarding the dilution effect

In this section we show how to compute an upper bound to the reduction in the expected number of false negatives obtained when implementing ordered pooling as opposed to random pooling using only partial information regarding the dilution effect. For a given pool size K , we show that only five parameters are required to compute the upper bound: μ , b , S_e , S_p and $h(1, K)$.

To get a tight upper bound to the maximum reduction in the expected number of false negatives obtained when implementing ordered pooling, we need to make the following assumption:

Assumption D.5

$$\mu(1 - \mu)^{K-2} \leq b(1 - b)^{K-2}. \quad (31)$$

While assumption D.5 is not very intuitive, notice that a sufficient condition for it to hold is that $b \leq \frac{1}{K-1}$. Indeed, noticing that the function $f(q) \equiv q(1 - q)^{K-2}$ is unimodal within the interval $[0, 1]$, having a single peak at $\frac{1}{K-1}$, we have that, if $b \leq \frac{1}{K-1}$, then b and μ both lie on the increasing side of the function $f(\cdot)$, in which case $\mu(1 - \mu)^{K-1} \leq b(1 - b)^{K-1}$, since $\mu \leq b$.

To assess whether assumption D.5 is realistic or not, one only needs to find the maximum $b > \mu$ such that inequality (31) holds and check whether or not it surpasses the actual b observed in the data. Figure IV depicts the cutoff points for b below which inequality (31) holds for different values of μ and K . As it can be seen from the figure, the inequality is more likely to hold the smaller the pool size, and the lower the overall prevalence rate μ .

Approximating $\bar{q}$ by μ (a good approximation when N is large) allows us to derive a simple formula to an upper bound to $\frac{FN_r^K(\mu, n, K) - FN_o^K(q_1, q_2, \dots, q_N)}{N}$.

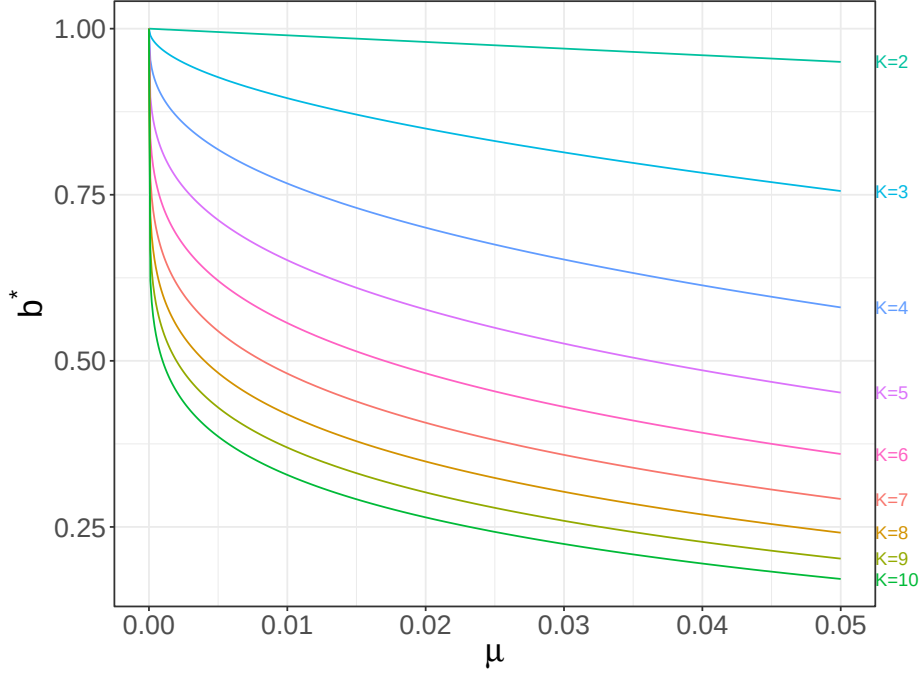


Figure IV: The $b^* \in (\mu, 1]$ such that $\mu(1-\mu)^{K-2} = b^*(1-b^*)^{K-2}$ for different combinations of μ and K .

Proposition D.6 *Upper bound to the reduction in the expected number of false negatives:* If assumptions 1, D.5 and D.1 hold, then, for a given realization of $(q_1, q_2, \dots, q_N)$ with $\sum_i q_i/N = \mu$, we have that

$$\frac{FN_r^K(\mu, n, K) - FN_o^K(q_1, q_2, \dots, q_N)}{N} \leq \mu S_e(S_e - h(1, K)) [(1-\mu)^{K-1} - (1-b)^{K-1}].$$

Proof: Clearly, the probability that an infected subject is incorrectly classified as not infected is decreasing in the probability of infection of the subjects with whom she is matched with. Therefore, the following expression corresponds to a lower bound to the total expected number of false negatives obtained when implementing ordered pooling with pools of size K :

$$\begin{aligned} \underline{FN}_o^K &= \sum_{i=1}^N q_i \left[\sum_{I=0}^{K-1} (1 - S_e h(I+1, K)) \binom{K-1}{I} b^I (1-b)^{K-1-I} \right] \\ &= N \bar{q} \left[\sum_{I=0}^{K-1} (1 - S_e h(I+1, K)) \binom{K-1}{I} b^I (1-b)^{K-1-I} \right] \\ &= N \mu \left[\sum_{I=0}^{K-1} (1 - S_e h(I+1, K)) \binom{K-1}{I} b^I (1-b)^{K-1-I} \right]. \end{aligned}$$

We wish to find an upper bound to

$$\frac{FN_r^K - FN_o^K}{N} = \mu \left[\sum_{I=0}^{K-1} (1 - S_e h(I+1, K)) \binom{K-1}{I} (\mu^I (1-\mu)^{K-1-I} - b^I (1-b)^{K-1-I}) \right].$$

Notice that

$$\begin{aligned} \frac{FN_r^K - FN_o^K}{N} &= \mu [(1 - S_e h(1, K)) ((1-\mu)^{K-1} - (1-b)^{K-1}) + \\ &\quad + \sum_{I=1}^{K-1} (1 - S_e h(I+1, K)) \binom{K-1}{I} (\mu^I (1-\mu)^{K-1-I} - b^I (1-b)^{K-1-I})]. \end{aligned} \quad (32)$$

Now notice that the sign of the second term inside the brackets from equation (32),

$$\sum_{I=1}^{K-1} (1 - S_e h(I+1, K)) \binom{K-1}{I} (\mu^I (1-\mu)^{K-1-I} - b^I (1-b)^{K-1-I}), \quad (33)$$

is ambiguous, as $\mu^I (1-\mu)^{K-1-I} - b^I (1-b)^{K-1-I}$ could, in principle, be either positive or negative. Because $\mu \leq b$, we have that $\mu^I (1-\mu)^{K-1-I} - b^I (1-b)^{K-1-I} \leq 0$ for all $I \geq 1$ if and only if $\mu(1-\mu)^{K-2} - b(1-b)^{K-2} \leq 0$, which, from assumption (D.5), holds. So, for every $I \geq 1$, the term $h(I, K)$ from expression (30) is always being multiplied by a non-positive term, which implies that this expression is increasing in $h(I+1, K)$ for all $I \in \{1, 2, \dots, K-1\}$. But from assumption D.1 we have that $h(I+1, K) \leq S_e$. Therefore, expression (33) is bounded from above by

$$\begin{aligned} &\sum_{I=1}^{K-1} (1 - S_e^2) \binom{K-1}{I} (\mu^I (1-\mu)^{K-1-I} - b^I (1-b)^{K-1-I}) = \\ &\quad - (1 - S_e^2) [(1-\mu)^{K-1} - (1-b)^{K-1}]. \end{aligned}$$

This implies that

$$\begin{aligned} \frac{FN_r^K - FN_o^K}{N} &\leq \mu [(1 - S_e h(1, K)) ((1-\mu)^{K-1} - (1-b)^{K-1}) \\ &\quad - (1 - S_e^2) [(1-\mu)^{K-1} - (1-b)^{K-1}]] \\ &\leq \mu S_e (S_e - h(1, K)) [(1-\mu)^{K-1} - (1-b)^{K-1}]. \end{aligned}$$

■

Notice that, in the absence of dilution effects, i.e., when $h(I, K) = S_e$ for all $I \geq 1$, the upper bound from proposition D.6 equals zero. This happens because, in the absence

of dilution effects, the probability that someone receives a false negative result does not depend on who the subject is matched with (e.g., see Aprahamian, Bish and Bish [2019]). In general, the larger the dilution effect (and therefore the lower $h(1, K)$ is), the greater is the reduction in false negatives that can be achieved by implementing ordered pooling as opposed to random pooling.

Remark D.7 *The upper bound to the reduction in the expected number of false negatives from proposition D.6 is decreasing in $h(1, K)$.*

So, if the tester does not have an accurate estimate of the dilution function $h(I, K)$, he can still conservatively estimate the upper bound from proposition D.6 by assuming a $h(1, K)$ well below S_e .

D.4 Numerical exercises with looser upper bounds

Figure V replicates the numerical exercise conducted in section 9, but this time using the looser upper bounds derived in this section. Similarly, figure VI replicates the numerical exercise conducted in section 10. In both cases, the upper bounds are much looser than the ones derived on the main text. But one can easily verify that, for small prevalence levels of a disease and small pool sizes, even these loose upper bounds can be very close to zero.

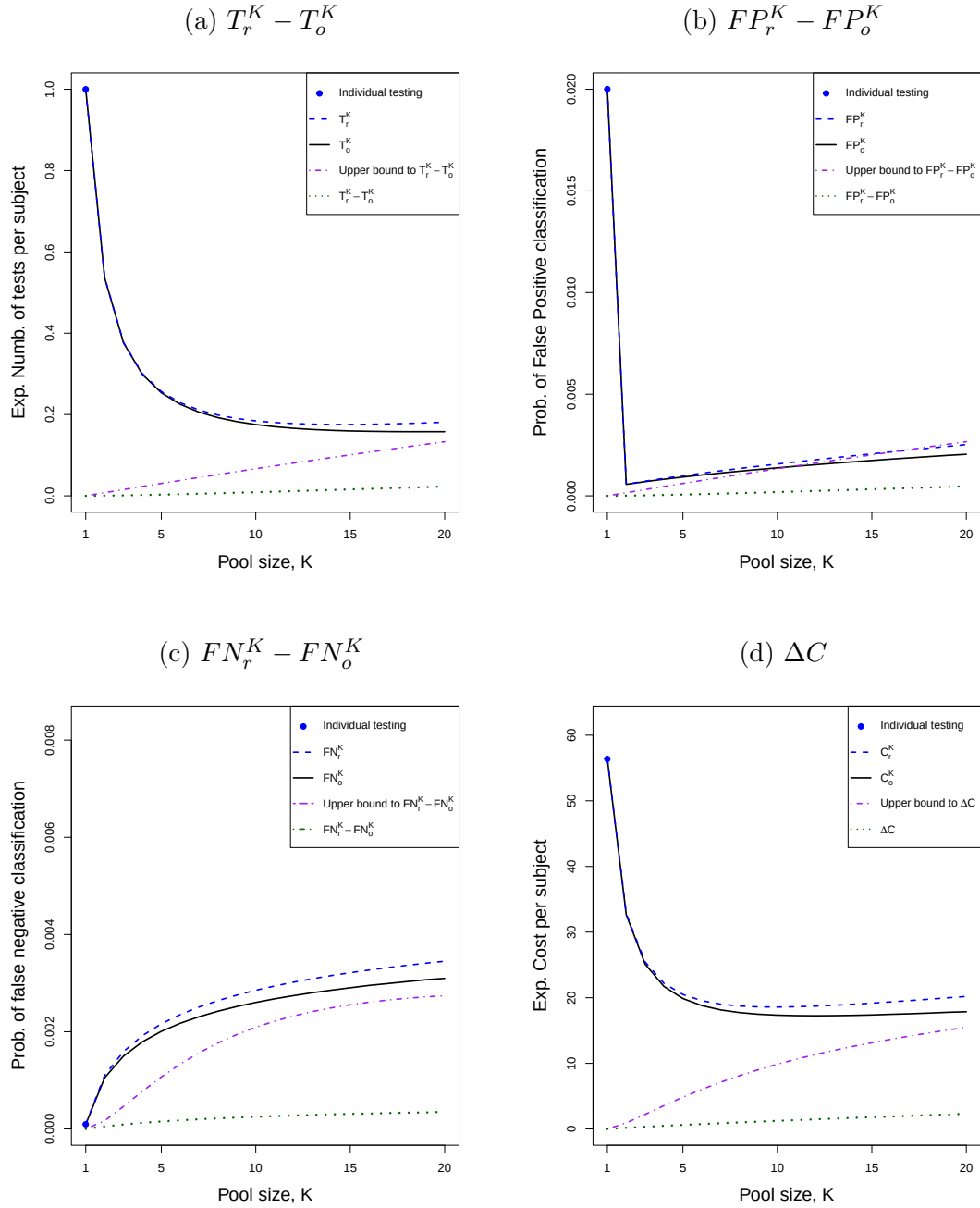


Figure V: Expected number of tests and classification errors per subject for Chlamydia, assuming $N = 300$. The blue dot corresponds to the case in which subjects are individually tested, i.e. $K = 1$. Parameters for the dilution function (16): $S_e = .99$, $S_p = .98$ and $\delta = 0.15$. The purple lines correspond to the looser upper bounds derived in section D.

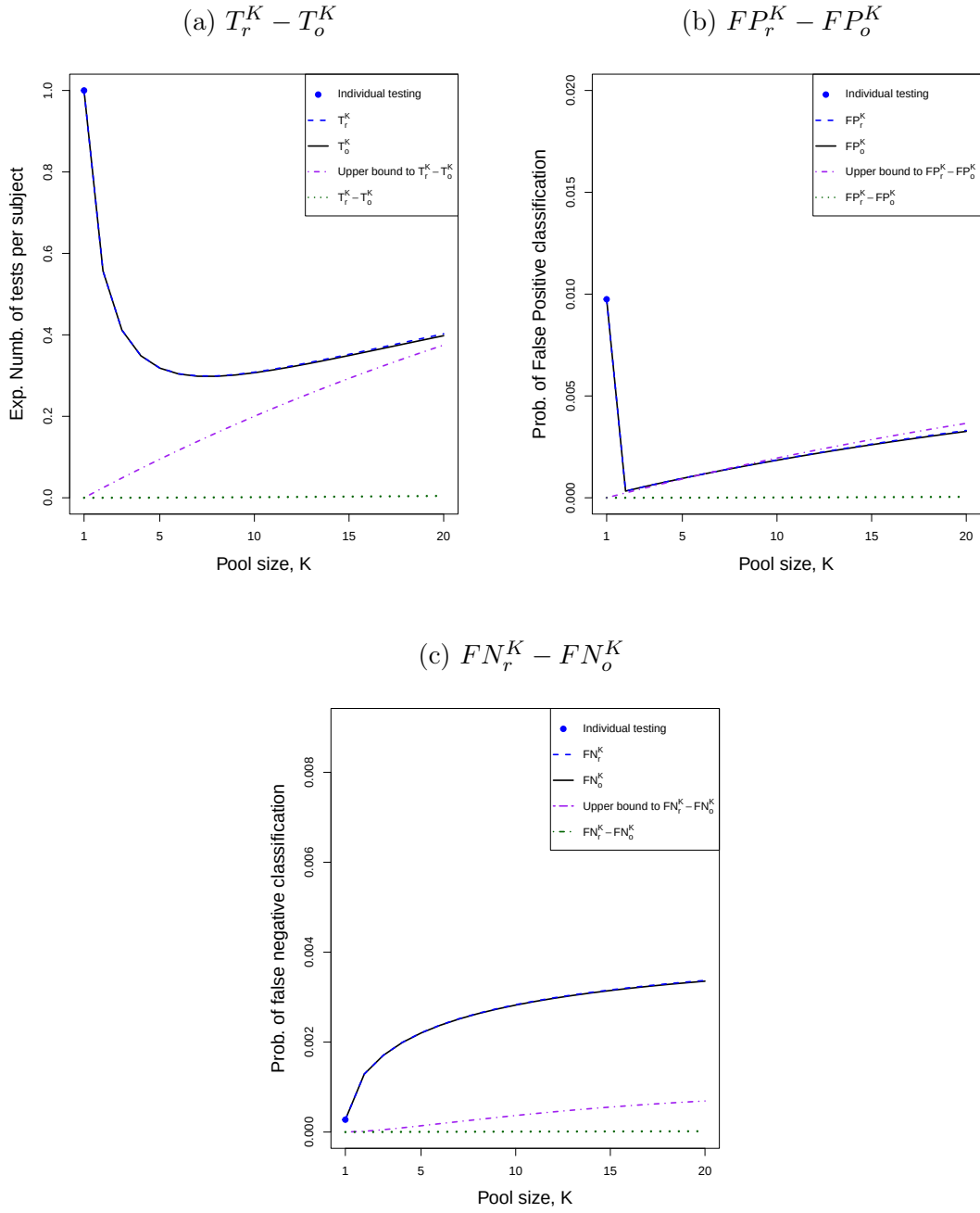


Figure VI: Expected number of tests and classification errors per subject for Covid-19, assuming $N = 300$. The blue dot corresponds to the case in which subjects are individually tested, i.e. $K = 1$. Parameters for the dilution function (16): $S_e = .99$, $S_p = .989$ and $\delta = 0.0459$. The purple lines correspond to the looser upper bounds derived in section D.

Bibliography

- Agoti, Charles N., Martin Mutunga, Arnold W. Lambisia, et al.** 2021. “Pooled testing conserves SARS-CoV-2 laboratory resources and improves test turn-around time: experience on the Kenyan Coast.” *Wellcome Open Research*, 5: 186.
- Aprahamian, Hrayr, Douglas R. Bish, and Erub K. Bish.** 2019. “Optimal Risk-Based Group Testing.” *Management Science*, 65(9): 4365–4384.
- Aprahamian, Hrayr, Ebru K. Bish, and Douglas R. Bish.** 2018. “Adaptive risk-based pooling in public health screening.” *IIE Transactions*, 50(9): 753–766.
- Basso, Leonardo J., Marcel Goic, Marcelo Olivares, et al.** 2023. “Analytics Saves Lives During the COVID-19 Crisis in Chile.” *INFORMS Journal on Applied Analytics*, 53(1): 9–31.
- Dorfman, Robert.** 1943. “The Detection of Defective Members of Large Populations.” *The Annals of Mathematical Statistics*, 14: 436–440.
- Saraiva, Gustavo Quinderé.** 2023. “Pool testing with dilution effects and heterogeneous priors.” *Health Care Management Science*, 651–672.