

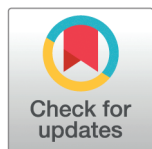
METHODS

A Bayesian model for repeated cross-sectional epidemic prevalence survey data

Nicholas Steyn^{1*}, Marc Chadeau-Hyam^{2,3}, Paul Elliott^{2,3}, Christl A. Donnelly^{1,4}

1 Department of Statistics, University of Oxford, Oxford, United Kingdom, **2** MRC Centre for Environment and Health, School of Public Health, Imperial College London, London, United Kingdom, **3** Department of Epidemiology and Biostatistics, School of Public Health, Imperial College London, London, United Kingdom, **4** Pandemic Sciences Institute, University of Oxford, Oxford, United Kingdom

* nicholas.steyn@univ.ox.ac.uk



Abstract

Epidemic prevalence surveys monitor the spread of an infectious disease by regularly testing representative samples of a population for infection. State-of-the-art Bayesian approaches for analysing epidemic survey data were constructed independently and under pressure during the COVID-19 pandemic. In this paper, we compare two existing approaches (one leveraging Bayesian P-splines and the other approximate Gaussian processes) with a novel approach (leveraging a random walk and fit using sequential Monte Carlo) for smoothing and performing inference on epidemic survey data. We use our simpler approach to investigate the impact of survey design and underlying epidemic dynamics on the quality of estimates. We then incorporate these considerations into the existing approaches and compare all three on simulated data and on real-world data from the SARS-CoV-2 REACT-1 prevalence study in England. All three approaches, once appropriate considerations are made, produce similar estimates of infection prevalence; however, estimates of the growth rate and instantaneous reproduction number are more sensitive to underlying assumptions. Interactive notebooks applying all three approaches are also provided alongside recommendations on hyperparameter selection and other practical guidance, with some cases resulting in orders-of-magnitude faster runtime.

OPEN ACCESS

Citation: Steyn N, Chadeau-Hyam M, Elliott P, Donnelly CA (2025) A Bayesian model for repeated cross-sectional epidemic prevalence survey data. *PLoS Comput Biol* 21(10): e1013515.
<https://doi.org/10.1371/journal.pcbi.1013515>

Editor: Sasikiran Kandula, Norwegian Institute of Public Health: Folkehelseinstituttet, NORWAY

Received: April 16, 2025

Accepted: September 10, 2025

Published: October 3, 2025

Copyright: © 2025 Steyn et al. This is an open access article distributed under the terms of the [Creative Commons Attribution License](https://creativecommons.org/licenses/by/4.0/), which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

Data availability statement: All source code and data (simulated and from the REACT-1 survey) used in this study are available at: <https://github.com/nicsteyn2/EpidemicSurveySmoothing>.

Funding: N.S. acknowledges support from the Oxford-Radcliffe Scholarship from University College, Oxford, the EPSRC CDT in Modern Statistics and Statistical Machine Learning (Imperial College London and University of Oxford). M.C-H. acknowledges support from Cancer Research UK, Population Research

Author summary

Understanding how infections spread in a population is crucial during an epidemic, and large-scale surveys that test people for current infection can provide valuable insights. These surveys are resource-intensive and the data they produce can be noisy and hard to interpret. In this study, we investigate how three different statistical approaches (two established and one novel) can explore such data. We found that some common modelling choices, particularly the treatment of observation noise, can meaningfully shape

Committee Project grant “Mechanomics” (grant 22184), the H2020-EXPANSE (Horizon 2020 grant 874627) and H2020-LongITools (Horizon 2020 grant 874739). P.E. is director of the MRC Centre for Environment and Health (MR/L01341X/1 and MR/S019669/1). P.E. acknowledges support from the Department of Health & Social Care in England (NIHR and UKRI, REACT-LC, COV-LT-0040), the Medical Research Council (MR/V030841/1), the MRC Centre for Environment and Health (MR/S019669/1), and the UK Health Security Agency for the REACT-3 study (2024-2025). C.A.D. acknowledges support from the MRC Centre for Global Infectious Disease Analysis, the NIHR Health Protection Research Unit in Emerging and Zoonotic Infections, the NIHR funded Vaccine Efficacy Evaluation for Priority Emerging Diseases (PR-OD-1017-20007), and the Oxford Martin Programme in Digital Pandemic Preparedness. The funders had no role in study design, data collection and analysis, decision to publish, or manuscript preparation.

Competing interests: We have read the journal's policy and the authors of this manuscript have the following competing interests: MCH holds shares in the OSMOSE consulting company. Consultancies are independent of the present work.

results. Our findings highlight the need for careful, robust methods to help researchers and public health officials make best use of existing data, design more effective surveys, and extract clearer insights from future studies.

Introduction

The aims of infectious disease surveillance include describing disease burden, monitoring trends, and identifying outbreaks and novel pathogens through the collection, analysis, and interpretation of health data [1]. A range of passive systems, such as automated reporting from healthcare facilities, and active systems, such as field investigations, are used to gather these data. Recently, novel surveillance methods, such as wastewater testing [2], mobile phone data [3,4], and social media monitoring [5], have been developed to complement traditional surveillance methods. While these methods provide an unprecedented volume of data, they are often subject to biases and limitations that make reliable inference and interpretation difficult [6].

Large-scale prevalence surveys are another tool for the surveillance of infectious diseases. These surveys typically use random sampling methods to produce estimates of the prevalence of infection in a given population. The quality of data collected also allows for the robust estimation of epidemiological quantities such as prevalence P_t , the growth rate r_t , and the instantaneous reproduction number R_t [7]. Established during the COVID-19 pandemic, the REal-time Assessment of Community Transmission (REACT-1) study in England [8] processed over 2.5 million self-administered throat and nose samples between May 2020 and March 2022. The ONS Coronavirus Infection Survey [9] monitored the spread of SARS-CoV-2 in the United Kingdom, also processing millions of samples over the course of the pandemic. These surveys have been instrumental in understanding, for example, the spread of SARS-CoV-2 [8], the dynamics of infection hospitalisation and infection fatality ratios [10,11], and the impact of vaccination [12].

Implementing such large-scale surveys is expensive; for example, the ONS Coronavirus Infection Survey cost £988.5 million (until September 2023) [13]. It is therefore critical to maximise the information extracted from the data - not only to justify cost, but also to support timely decision-making and improve policy relevance. This is made challenging by the substantial noise inherent in point prevalence estimates. Even with an ideal survey design, large sample size, and only ignorable non-response, the coefficient of variation of common binomial proportion estimators can be large [14], particularly when infection prevalence is low.

Even when the goal is to perform “model-free” inference (making inferences that reflect only the data, rather than modelling assumptions), noise in these data often necessitates the use of smoothing methods to improve the quality of the outputs. Smoothing methods impose constraints on day-to-day variability, allowing data from multiple days to inform point estimates. Smoothing allows the researcher to make more confident statements, increasing information yield or reducing the sample size needed for a given confidence level. However, all smoothing methods make assumptions about the underlying process, whether explicitly or implicitly. Even seemingly simple smoothing methods can introduce substantial bias, turning “model-light” estimates (that rely only minimally on modelling assumptions) into “model-heavy” estimates (that are strongly shaped by modelling assumptions) that may not accurately reflect reality.

Here we introduce a novel Bayesian approach for smoothing and performing inference on epidemic survey data, referred to as SIMPLE (Survey Inference Method for Prevalence

and other Latent variables in Epidemiology). This approach is based on hidden-state models and sequential Monte Carlo (SMC) methods [15,16], and is designed to be flexible enough to incorporate key assumptions of existing approaches, while avoiding some of the computational and mathematical complexity of these approaches. We compare the SIMPLE approach to two existing approaches, one by Eales et al. [17] that leverages Bayesian P-splines and another by Abbott and Funk [7] that uses Gaussian processes approximations. We refer to these approaches as the Eales approach and the Abbott approach respectively. All three approaches are presented using common and general notation, allowing for direct comparison of their assumptions, performance, and results.

The results of this paper are structured in three parts. First, we use the SIMPLE approach to investigate the impact of survey design (such as the number of samples and individual response bias) and underlying epidemic dynamics (such as the variability in the growth rate) on the quality of estimates, highlighting key modelling decisions that should be made when analysing epidemic survey data. Second, we demonstrate how the approaches of Eales and Abbott can be adapted to account for these considerations, and compare all three (improved) approaches on simulated data to highlight similarities and differences in their performance. The use of simulated data provides a ground truth that allows us to calculate and compare quantities such as statistical coverage (the percentage of times that credible intervals of a given level contain the true value). Third, we apply all three approaches to real-world data from the REACT-1 study, covering the COVID-19 pandemic between May 2020 and March 2022 in England, comparing the quality of estimates and computational requirements of each approach. We provide recommendations for future modelling and survey design based on our findings.

Our goal is to provide a framework for understanding the impact of modelling decisions on the quality of estimates, to consolidate and improve state-of-the-art methods, and to make recommendations for future modelling and survey design. We also provide documented source code for all three approaches, notebooks demonstrating their use, and recommendations for hyperparameter selection and other practical considerations.

Materials and methods

The SIMPLE approach

We introduce a suite of state-space models for simultaneously smoothing and performing inference on epidemic survey data. Each state-space model consists of an explicitly defined epidemic and observation model. The epidemic model encodes our assumptions about the unobserved dynamics of the epidemic, while the observation model describes how the observed data are generated from the underlying epidemic.

Two epidemic models are considered: one for estimating the daily exponential growth rate r_t in prevalence P_t and one for estimating the instantaneous reproduction number R_t , infection incidence I_t and prevalence P_t . Three observation models are considered: a basic model that assumes the number of positive swabs (diagnostic tests) follows a binomial distribution (as used by the original Eales approach), an extra-binomial model that accounts for overdispersion in the data, and a weighted model that accounts for survey weights.

We refer to the unknown time-varying quantities of interest, such as r_t and P_t , as hidden states. Other unknown parameters that are not time-varying, such as the level of overdispersion in the extra-binomial model, are referred to as static parameters.

Throughout this paper, we use the term *prevalence* to describe the proportion of the population that would test positive. This quantity may differ from the proportion of the population

with an active infection (due to imperfect test sensitivity/specificity) and from the proportion who are infectious (as test positivity can outlast infectiousness).

Growth rate epidemic model. This model assumes that the daily growth rate r_t in prevalence follows a Gaussian random walk, encoding an assumption that r_t varies smoothly over time. Prevalence P_t on day $t \in \{1, 2, \dots, T\}$ then varies exponentially with rate r_t .

$$r_t = r_{t-1} + \epsilon_t, \quad \epsilon_t \sim N(0, \sigma^2) \quad (1)$$

$$P_t = P_{t-1} e^{r_t} \quad (2)$$

Coupled with initial distributions for r_0 and P_0 , and a prior distribution for σ , Eqs 1 and 2 define the entire epidemic model. Parameter σ controls the smoothness of our estimates and is inferred from the data. As this model makes minimal assumptions about the underlying epidemic, we use it as the default epidemic model unless otherwise stated.

We use default initial distributions of $r_0 \sim \text{Uniform}(-0.3, 0.3)$ and $P_0 \sim N(\hat{p}_1, \hat{p}_1(1 - \hat{p}_1)/n_1)$, where $\hat{p}_1 = n_1^+/n_1$ (the Wald interval), and default prior distribution $\sigma \sim \text{Uniform}(0, 0.2)$, although inferences are insensitive to the choice of these prior distributions (Sect A of S1 Text).

Taking a function space view, modelling the growth rate r_t as a first-order random walk implies a log-prior probability proportional to $-\frac{1}{2\sigma^2} \sum (r_t - r_{t-1})^2$. This can be viewed as a discrete approximation of the integral penalty $\int \left(\frac{d}{dt} r_t\right)^2 dt$, thus implying that the equivalent continuous growth rate function lies in the $W^{1,2}$ Sobolev space (the space of functions in L^2 that have square-integrable weak first derivatives). Although these are locally smooth functions, they do not allow for abrupt changes in r_t which may arise from sudden changes in behaviour or policy [18].

Reproduction number epidemic model. Alternatively, we may want to estimate the instantaneous reproduction number R_t , a popular alternative measure to r_t for characterising the rate of epidemic spread [19]. R_t is the average number of secondary infections generated by a primary infection at time t if an individual were to undergo their entire infectious period at this time. As with r_t , we employ a Gaussian random walk to smooth estimates, now on $\log R_t$ to ensure positivity:

$$\log R_t = \log R_{t-1} + \epsilon_t, \quad \epsilon_t \sim N(0, \sigma_R^2). \quad (3)$$

We then employ the renewal model [16,20], which relates past infection incidence $I_{1:t-1}$ to current infection incidence I_t through the instantaneous reproduction number R_t and a generation time distribution (representing the time from infection of an infector and their infectee, described by probability mass function w_u):

$$I_t = R_t \sum_{u=1}^t I_{t-u} w_u. \quad (4)$$

Finally, we relate prevalence P_t to infection incidence I_t through a test-sensitivity function h_u that describes how likely an individual is to test positive u days after infection. We do not consider imperfect test specificity, although this could be included with the addition of a constant term to P_t :

$$P_t = \sum_{u=1}^t I_{t-u} h_u. \quad (5)$$

When applying the reproduction number epidemic model to data from the REACT-1 study, the generation time is assumed to follow a gamma distribution with mean 6.4 days and standard deviation 4.2 days [7], chosen to be consistent with the generation time distribution used in the published Abbott approach. For simplicity, we discretised this distribution by evaluating the gamma density function at integer time steps and normalising so $\sum w_u = 1$, although other discretisation methods can be more accurate [21,22]. There is also evidence that the generation time of SARS-CoV-2 has shortened with more recent variants [23], and this should be kept in mind when interpreting our real-world results. In the absence of a REACT-1-specific test-sensitivity function, we use the pointwise central estimate of h_u from [24], reflecting the type of test used in REACT-1 (reverse transcription polymerase chain reaction (RT-PCR)), but not other study-specific factors that may impact sensitivity. This is the same mean test-sensitivity function as used in the Abbott approach.

Estimating R_t requires many additional assumptions about the underlying epidemic, and thus has more potential to bias estimates of P_t . We focus on the growth rate epidemic model for much of this paper, although demonstrate the R_t estimator on real-world REACT-1 data.

Basic observation model. Given n_t swabs conducted on day t , the observed number of positive swabs n_t^+ is modelled as a binomial random variable with probability P_t :

$$n_t^+ \sim \text{Binomial}(n_t, P_t). \quad (6)$$

This is also the observation model used in the Eales approach as originally published.

Extra-binomial observation model. Prevalence data are often overdispersed relative to the binomial distribution, exhibiting what is known as extra-binomial variation [25]. If we assume that the “observable” prevalence at time t is beta-distributed with mean equal to prevalence P_t and variance $\rho P_t(1 - P_t)$, our observation distribution becomes:

$$n_t^+ \sim \text{Beta-binomial}(n_t, \alpha_t, \beta_t), \quad (7)$$

where $\alpha_t = P_t \left(\frac{1}{\rho} - 1 \right)$ and $\beta_t = (1 - P_t) \left(\frac{1}{\rho} - 1 \right)$.

The additional parameter $\rho \in (0, 1)$ controls the modelled level of overdispersion in the data. Larger values of ρ indicate greater levels of overdispersion while letting $\rho \rightarrow 0$ recovers the binomial model. By default, we use a Uniform(0, 0.01) prior distribution for this parameter. For context, the greatest upper bound of any 95% credible interval for ρ estimated from a real-world dataset is 0.0006.

This is our default observation model and is used in all subsequent analyses unless otherwise stated. It is common in other epidemiological settings to model overdispersed count data using the negative-binomial distribution [26], which is useful when the data have no natural upper limit. However, as the number of positive swabs are bounded above by the total number of swabs, we do not consider this here.

Weighted observation model. Survey weights are used to account for bias in survey data arising from unequal probabilities of selection and/or response. Letting $w_{t,i}$ be the (normalised) weight assigned to the i^{th} sample taken at time t , and $x_{t,i}$ be the corresponding result (where $x_{t,i} = 1$ if individual i on day t tests positive, and zero otherwise), the weighted swab positivity is:

$$\hat{p}_t = \sum_{i=1}^{n_t} x_{t,i} w_{t,i}. \quad (8)$$

If the weights are uncorrelated with swab positivity, then $\text{Var}(\hat{p}_t) = P_t(1 - P_t) \sum_{i=1}^{n_t} w_i^2$. However, if the weights are correlated with the individual probability of testing positive, the variance of $\hat{p}_t$ may be over or underestimated by this expression. The difference depends on the specific relationship between weights and outcome, which is generally unknown. Thus, we model $\hat{p}_t$ using a normal distribution with mean P_t and variance $cP_t(1 - P_t) \sum_{i=1}^{n_t} w_{t,i}^2$, where c is an estimated scalar parameter:

$$\hat{p}_t \sim N\left(P_t, cP_t(1 - P_t) \sum_{i=1}^{n_t} w_{t,i}^2\right). \quad (9)$$

The observation distribution is no longer exact, in that we are approximating the unknown distribution of $\hat{p}_t$ with a normal distribution. This is similar to the approximation used in the Abbott approach. By default, we use a Uniform(0.1, 10) prior distribution for c .

Sequential Monte Carlo (SMC). We use an SMC algorithm (also known as a particle filter), the bootstrap filter, to estimate the posterior distributions of the hidden states (such as r_t , P_t , and R_t) at each time step t , given the observed data [16]. We also use this algorithm to estimate the posterior predictive distribution of n_t^+ , which is similar to the frequentist's predictive distribution, and can be used to assess the quality of fit of the model when fitting to real-world data. Particle marginal Metropolis-Hastings (PMMH) is used to estimate the static parameters, the uncertainty of which is then marginalised over in the final inference.

We run three chains of the PMMH algorithm, each for 100 iterations at a time, until the maximum Gelman-Rubin diagnostic ($\hat{R}$) [27] is less than 1.05 and the minimum effective sample size (ESS) is greater than 100. These algorithms are detailed in full in Sect B of [S1 Text](#). We make the full source code (written natively in Julia [28]) available on [GitHub](#), along with [notebooks](#), to facilitate the application of these approaches to other datasets.

When fitting the model for the reproduction number, an additional “wind-in” period (default 10 days) is required to account for infection history prior to the first observation. After sampling from the standard prior distribution for P_0 , assumed values of $L_{9:0}$ are set to $P_0 / \sum h_u$, ensuring the distribution of implied prevalence at time 0 is consistent with the chosen prior distribution. As 10 days may truncate the generation time distribution and test-sensitivity function, we rescale the generation time distribution to sum to 1 and the test-sensitivity function to sum to its original total over the combined length of past data and the wind-in period.

The Eales approach

The Eales approach [17] models logit-transformed P_t using Bayesian p-splines:

$$\text{logit}P_t = \sum_{i=1}^N b_i B_{i,n}(t), \quad \text{logit}(x) = \log \frac{x}{1-x}, \quad (10)$$

where $B_{i,n}(t)$ are basis functions and b_i are estimated spline coefficients. The coefficients give the value of the spline at the corresponding “knots” t_i while the basis functions allow for interpolation. By default, fourth-order basis functions $B_{i,4}(t)$ are used. We refer the reader to the original paper for the full construction [17].

The smoothness of these splines is controlled by the spacing of the spline knots (set by the user), and a second-order Gaussian random walk prior distribution on the spline coefficients,

with standard deviation σ_{Eales} (estimated from the data):

$$b_i - b_{i-1} = b_{i-1} - b_{i-2} + u_i, \quad u_i \sim N(0, \sigma_{Eales}^2). \quad (11)$$

As $\text{logit}P_t - \text{logit}P_{t-1} \approx \log P_t - \log P_{t-1} = r_t$ for small P_t , this is approximately equivalent to modelling the growth rate using a Gaussian random walk, a similar smoothing assumption to our growth rate epidemic model. This approximation can be improved by noting that $\text{logit}P_t - \text{logit}P_{t-1} \approx \log P_t - \log P_{t-1} + (P_t - P_{t-1})$ and that the smoothness of P_t implies that $P_t - P_{t-1}$ is small. Taking a function space view, by directly penalising the second-order differences of the spline coefficients, the Eales approach places logit-prevalence in the $W^{2,2}$ Sobolev space. Since the (continuous) growth rate is approximately the first derivative of logit-prevalence, the implied growth rate function lies in the $W^{1,2}$ Sobolev space, the same space implied by the SIMPLE approach.

By default, spline knots are placed approximately every 5 days (exactly 5 days when the duration of the data is divisible by 5) to balance flexibility with computational efficiency. As we generally work in integer time, there is no benefit to knot spacing shorter than 1 day. If knots are placed on each day, we only need to evaluate the splines at the knot locations, and thus can work with b_i directly (i.e., we do not need to employ any splines). In this case, except for modelling the change in $\text{logit}P_t$ instead of $\log P_t$, the model is equivalent to the SIMPLE (extra-binomial) model. We examine this equivalence further in Sect C of [S1 Text](#).

The original Eales approach modelled positive swabs n_t^+ with a binomial distribution, equivalent to our basic observation model. We update this to use a beta-binomial distribution with overdispersion parameter ρ , equivalent to our extra-binomial observation model, improving both convergence of the algorithm and the quality-of-fit of the model (Sects C and D of [S1 Text](#)). We do not provide a weighted implementation of the Eales approach, although this could be achieved by using the weighted observation model described above.

Growth rates (in prevalence) are back-calculated from the fitted splines by setting $r_t = \log(P_t/P_{t-1})$, where $P_{1:T}$ are sampled from the posterior distribution of the splines. A weakly informative inverse-gamma prior distribution ($\alpha = \beta = 0.0001$) is used for σ_{Eales} , a uniform prior distribution is used for ρ , and a uniform prior distribution is used for the first two spline coefficients.

Rather than explicitly model incidence, the Eales approach makes a series of simplifying assumptions when estimating the reproduction number. Specifically, it assumes that at each independent time step t , R_t has been fixed for the past τ days (a trailing-window approach akin to that used in EpiEstim [20]). If τ is chosen to be larger than both the maximum generation time and maximum duration of swab positivity, then the renewal equation (Eq 4) can be used to estimate R_t directly. In practice, $\tau = 14$ days is used to prevent oversmoothing of R_t , despite this likely being less than the maximum duration of test sensitivity for SARS-CoV-2 (which, according to [24], remains above 10% even after 21 days after infection). The trailing-window approach is known to produce biased estimates of R_t [29], which can be partially accounted for by reporting estimates shifted by $\tau/2$ days [30].

Eales et al. [17] provide source code to fit this model using the R programming language [31] and Rstan. We adapt their code to fit the model using the actively developed interface to Stan, cmdstanr [32]. As with Rstan, inference is performed using Hamiltonian Monte Carlo (HMC) with the adaptive No-U-Turn Sampler (NUTS). We also adjust the hyperparameters of their algorithm (increasing the maximum tree-depth to 15 and decreasing the number of warm-up/sampling iterations to 200/300) to substantially improve convergence times, from at least 30 hours on the full REACT-1 dataset to less than 1 hour (Sect D of [S1 Text](#)).

As in the SIMPLE approach, convergence is measured by running three chains and checking that maximum $\hat{R} < 1.05$ and minimum $ESS > 100$. In this case, “hidden states” such as the growth rate are estimated alongside static parameters, so also feature in the convergence checks. This means the maximum $\hat{R}$ and minimum ESS diagnostics are not directly comparable to those of the SIMPLE approach, although they do reflect the computationally expensive aspects of each approach (estimating the static parameters in the SIMPLE approach is much more expensive than estimating the hidden states).

The Abbott approach

The Abbott approach [7] models the first-order difference (other order differences are possible but not considered here) of logit-transformed daily infection incidence using a Laplace eigenfunction approximation [33,34] to a zero-mean Gaussian process:

$$\text{logit}(I_t) = \text{logit}(I_0) + \sum_{s=1}^t GP_s, \quad (12)$$

where I_0 is the initial incidence and GP_s is the value of the approximated Gaussian process at time s . At small incidence values I_t , this is approximately equivalent to modelling the growth rate in incidence r_t^{inc} as a Gaussian process:

$$r_t^{inc} = \log I_t - \log I_{t-1} \approx \text{logit} I_t - \text{logit} I_{t-1} = GP_t. \quad (13)$$

A squared exponential kernel, parameterised by variance σ_{Abbott}^2 and lengthscale ℓ , is used for the Gaussian process:

$$k(GP_t, GP_s) = \sigma_{Abbott}^2 \exp\left(-\frac{(t-s)^2}{2\ell^2}\right). \quad (14)$$

If the lengthscale ℓ is large relative to unit time steps, this Gaussian process can be locally approximated by an AR(1) process: $GP_{t+1} = \phi GP_t + \epsilon_t$ where $\phi = \exp(-1/2\ell^2)$ and $\epsilon_t \sim N(0, \sigma_{Abbott}^2(1 - \phi^2))$. This approximation highlights a mild similarity to our basic model: the model used in the Abbott approach can be (very loosely) viewed as a Gaussian random walk, except on r_t^{inc} and with a drift towards zero growth. The strength of this drift is controlled by ℓ with larger values of ℓ indicating a slower drift towards zero growth. From a function space perspective, the sample paths of a Gaussian process with a squared exponential kernel are almost surely infinitely differentiable [35]. Because prevalence is a convolution of incidence, the implied prevalence in this approach inherits the same very high degree of smoothness. While outside the scope of this paper, alternative Gaussian process kernels could be used in the Abbott approach to allow for only local smoothness (the Matérn kernel is a popular choice with direct links to the function spaces of the SIMPLE and Eales approaches [35]).

Prevalence P_t is modelled as a convolution of past incidence with a test-sensitivity function that describes how likely an individual is to test positive u days after infection (Eq 5). During the COVID-19 pandemic, estimates of the test-sensitivity function produced by Hellewell et al. [24] were used. We use the same estimates in our implementation (including in simulated data where applicable), although these curves depend on the pathogen (including variant, in the case of SARS-CoV-2 [36]), testing procedures (e.g. swabbing technique [37] and storage during transport [12]), and the population being tested (e.g. age structure and vaccination status [38]). As a result, these curves are not universal, and should be separately estimated for

each scenario. We compare multiple estimates of these curves, and the sensitivity of model results to these curves, in Sect E of [S1 Text](#).

The original Abbott approach used a normal likelihood for observed swab positivity with mean P_t and an estimated variance, a necessary simplification given the format of their data. We employ the beta-binomial observation model instead, although the binomial or weighted observation models are also possible and can be implemented by changing only a few lines of the Stan code. The original model also leveraged antibody data, accounting for vaccination status, in their model. We focus on swab positivity data so do not consider this here. Their antibody model could be included in any of the three approaches discussed in this paper. We leave this for future work.

Growth rates are back-calculated from the estimated prevalence by setting $r_t = \log(P_t/P_{t-1})$. A weakly informative inverse-gamma prior distribution is used for ℓ , a zero-mean normal distribution with standard deviation 0.1 is used as the prior distribution for σ_{Abbott} , a uniform prior distribution is used for ρ , a normal prior distribution is used for the initial logit-incidence (at the beginning of the 50-day wind-in period) with mean -4.6 and standard deviation 2, and a zero-mean normal prior distribution with standard deviation 0.25 is used for the initial growth rate in incidence (at the beginning of the 50-day wind-in period).

As the Abbott approach models infection incidence directly, the reproduction number can be estimated by directly applying the renewal model (Eq 4) to the estimated incidence curves.

Abbott and Funk provide source code in R and cmdstanr to fit their model. When fitting their model, we use 200 warm-up and 300 sampling iterations. As in the Eales approach, convergence is measured by running three chains and checking that maximum $\hat{R} < 1.05$ and minimum $ESS > 100$. This source code assumes data of a different format to ours (their inputs are daily means and credible intervals from a different Bayesian model), so we adapt their code to fit to raw data. We also remove code relating to their antibody model. The modified code is provided on GitHub.

Simulated data

Simulated datasets, approximately reflecting the dynamics of COVID-19 in England, are generated by sampling from the prior predictive distribution of a given model with predetermined parameter values. In general, T days (default 100) of an epidemic are simulated with a fixed initial prevalence P_0 (default 1%) and growth rate r_0 (default 0). As the Abbott approach models incidence and requires a wind-in period (default 50 days when simulating), we instead fix infection incidence at the start of the wind-in period at 0.1%. Other model-specific parameters are set to default values as outlined in [Table 1](#) unless otherwise stated. These are chosen to reflect estimates from the REACT-1 study, thus representing plausible values for SARS-CoV-2-like epidemics. To prevent unrealistic scenarios, we bound simulated prevalence in the range (0.1%, 30%) by resampling if a simulation exceeds these bounds. Example simulations are shown in [Fig 1](#).

Weighted survey data are generated by sampling artificial survey weights $w_{t,i}$ from a log-normal distribution (default log-mean 0 and scale $\xi = 0.5$) and normalising so $\sum_{i=1}^{n_t} w_{t,i} = 1$. To introduce bias, the probability that individual $i \in \{1, \dots, n_t\}$ tests positive is set to:

$$P_{t,i} = P_t \frac{w_{t,i}}{\sum_{j=1}^{n_t} w_{t,j}}. \quad (15)$$

This results in individuals with higher survey weights (i.e., individuals that are underrepresented in the sample) being more likely to test positive, while the weighted average of $P_{t,i}$ is

Table 1. Default parameter values used when simulating data, unless otherwise stated. These values are chosen to reflect estimates from fitting each model to the REACT-1 study. Prevalence and incidence are as a proportion of the population.

Parameter	Default value	Description
General parameters		
T	100 days	Duration of the simulated epidemic
P_0	0.01	Initial prevalence
r_0	0	Initial growth rate (per day)
$n_t, t \in \{1, \dots, T\}$	5,000	Daily swabs
SIMPLE approach (growth rate epidemic model)		
σ	0.01	Standard deviation of random walk on r_t (per day)
ρ	2×10^{-4}	Overdispersion parameter (beta-binomial observation distribution)
ξ	0.5	Scale of log-normal distribution for simulated weights
Eales approach		
σ_{Eales}	0.1	Standard deviation of random walk on spline knots (per knot)
ρ	2×10^{-4}	Overdispersion parameter (beta-binomial observation distribution)
Abbott approach		
σ_{Abbott}	0.08	Gaussian process kernel amplitude
ℓ	10	Gaussian process kernel lengthscale
ρ	2×10^{-4}	Overdispersion parameter (beta-binomial observation distribution)
I_{-init}	0.001	Initial incidence at the start of the wind-in period
r_{-init}	0	Initial growth rate in incidence (per day) at the start of the wind-in period

<https://doi.org/10.1371/journal.pcbi.1013515.t001>

equal to P_t . The simulated dataset then consists of $\{x_{t,i}, w_{t,i}\}$ where each $x_{t,i}$ is a realisation of a Bernoulli($P_{t,i}$) random variable. This imposes a linear relationship between survey weight and individual swab positivity, so represents a high level of bias for any assumed scale ξ . Default $\xi = 0.5$ was chosen to reflect the observed distribution of per-round survey weights in the REACT-1 study. An empirical analysis of these weights is provided in Sect F of S1 Text.

To consider the effect of survey design and epidemic dynamics, we fit each of the SIMPLE observation models (basic, extra-binomial, and weighted) to simulated datasets of duration $T = 100$ days from each observation model (with the growth rate epidemic model). Five values of the daily sample size n_t are considered (10, 100, 1,000, 10,000, 100,000) and two values of σ (0.008, 0.016). When simulating from the extra-binomial model we assume $\rho = 2 \times 10^{-4}$, and when simulating from the weighted model we assume $\xi = 0.5$. The choice of these values were guided by estimates from the REACT-1 study (see the REACT-1 results for ρ and Sect F of S1 Text for ξ). Each simulation is repeated 10 times to average over stochasticity. To evaluate model performance, we consider the coverage and average width of the 95% credible intervals for estimated prevalence P_t . Equivalent results for the growth rate in swab positivity r_t are presented in Sect G of S1 Text.

The REACT-1 study

The REACT-1 study was an infection prevalence survey that tested for SARS-CoV-2 infection in England between May 2020 and March 2022 [8]. Conducted over 19 rounds, a total of 2.5 million self-administered throat and nose samples were processed using RT-PCR. Daily swab positivity and sample sizes for all 19 rounds of the REACT-1 study are reported in Fig 2. An average of 6,236 samples were taken on each of the 400 days that sampling occurred, or an average of 3,564 samples per day over the 700 days spanned by the study (from 1 May 2020 to 31 March 2022 inclusive).

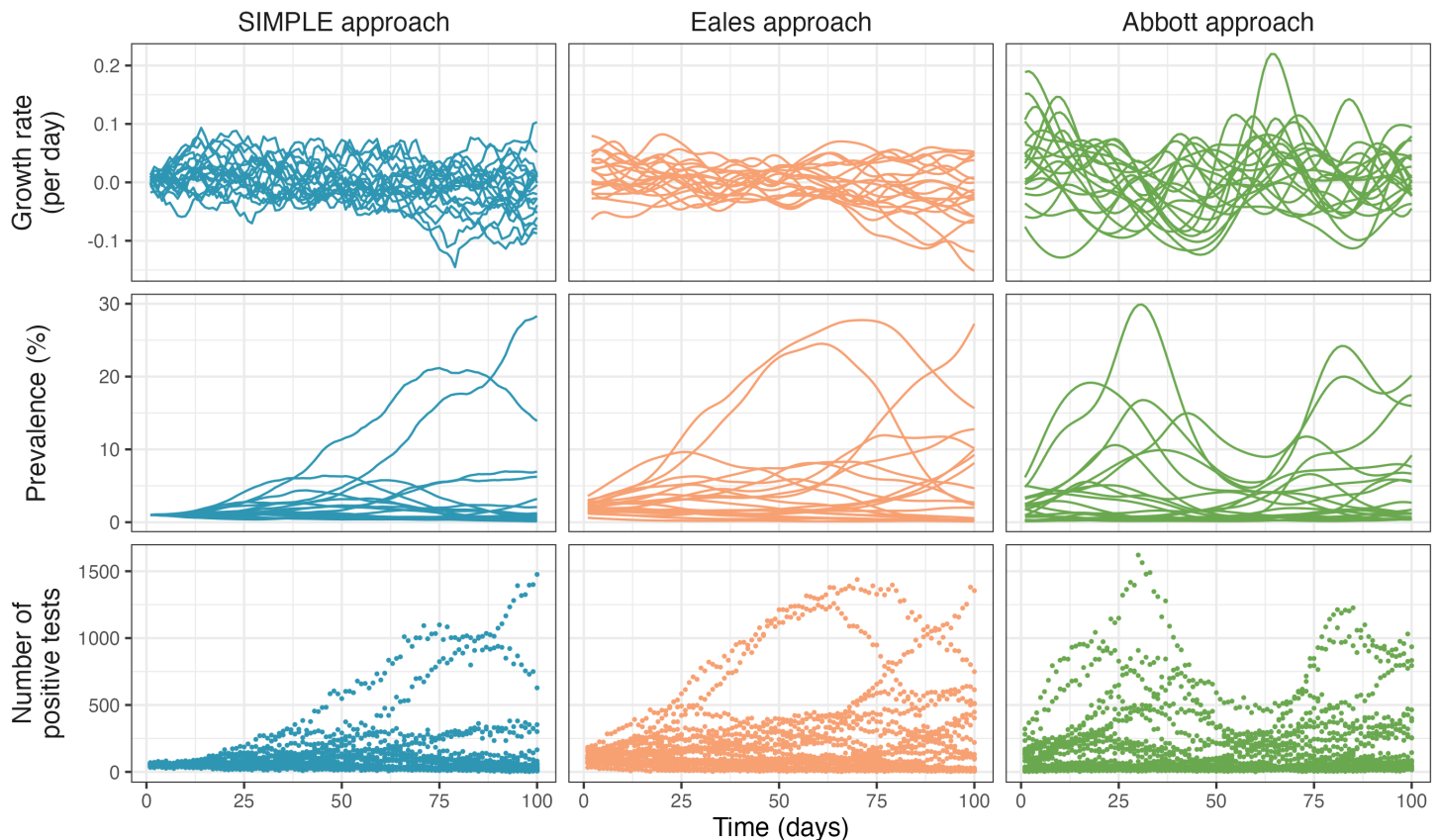


Fig 1. Simulated epidemic trajectories from the SIMPLE, Eales, and Abbott approaches. Default parameter values as described in Table 1 were used for these simulations. A total of 20 simulations are shown per approach, reflecting the range of possible outcomes. The top row shows the simulated growth rate r_t , the middle row shows the simulated prevalence P_t , and the bottom row shows the simulated number of positive swabs n_t^+ .

<https://doi.org/10.1371/journal.pcbi.1013515.g001>

Results

Survey design and epidemic dynamics

In this section we use the SIMPLE approach to investigate the impact of survey design, particularly the number of daily samples and individual response bias, and underlying epidemic dynamics (the variability in the growth rate) on the quality of estimates. Specifically, we consider when each of the three observation models (basic, extra-binomial, and weighted) are appropriate, and how well they perform when fit to simulated data from each model.

All three observation models produce well-calibrated and similarly narrow credible intervals when fit to simulated data from the basic observation model (i.e., a survey with binomial-distributed observations), despite the extra-binomial and weighted observation models featuring an additional (and in this case, unnecessary) parameter (Fig 3, column A). The width of these credible intervals decreases as n_t increases, with 1000 daily samples generally sufficient to produce credible intervals less than 1 percentage point in width, although this is highly simulation-dependent. Larger sample sizes may also allow for the estimation of prevalence by region and/or demographic, as seen in the REACT-1 study. When fit to simulated data with more variable growth rates (higher σ , dashed lines), the width of the credible intervals increases slightly, suggesting that a minor increase in daily sample size may be required when the growth rate is changing faster (to achieve the same level of precision). Finally, the

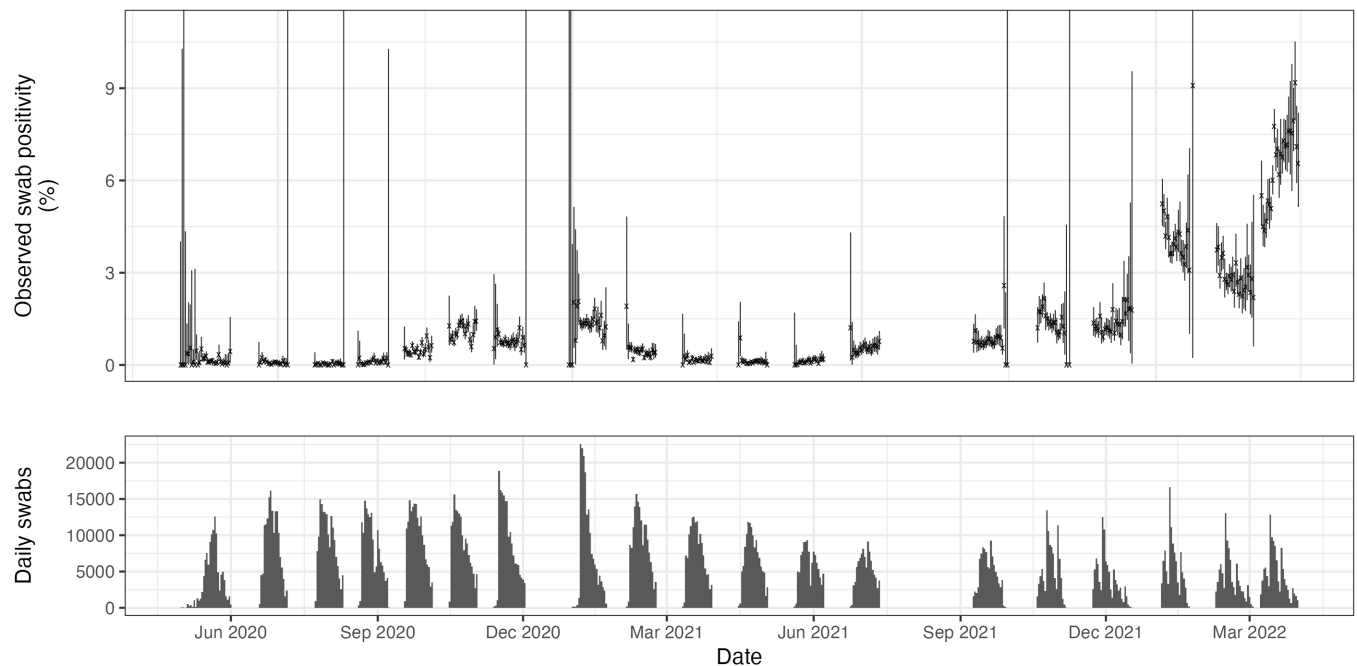


Fig 2. Daily SARS-CoV-2 swab positivity in England from the REACT-1 survey (upper) and corresponding daily sample sizes (lower). Daily 95% confidence intervals (vertical lines) for prevalence were calculated using the Agresti-Coull method [14] in the *binconf* function of the *Hmisc* package in R [39].

<https://doi.org/10.1371/journal.pcbi.1013515.g002>

weighted observation model produces poor coverage at low sample sizes, due to the breakdown of the normal approximation to the binomial distribution.

When fit to simulated data featuring extra-binomial variation, only the extra-binomial model consistently produces well-calibrated estimates of P_t (Fig 3, column B). At larger values of n_t , when the normal approximation to the binomial distribution is valid, the weighted model also produces well-calibrated estimates of P_t , as the additional parameter c allows the model to account for the extra-binomial variation. The basic model, while well calibrated at smaller values of n_t , produces poorly calibrated estimates as n_t increases, an example of simple modelling assumptions leading to “model-heavy” inferences. This is due to the basic model assuming that the observation variance is inversely proportional to n_t , forcing the credible intervals on P_t to decrease in width as n_t increases, even if additional observation noise is present. The basic model initially accounts for this by artificially increasing the estimated value of σ , allowing variation in r_t to capture the additional variation in n_t^+ (at the cost of biased estimates of σ and r_t - see Sect G of S1 Text), although this breaks down as n_t gets very large. This has real-world implications as seen on the REACT-1 dataset, where the original Eales approach (using a binomial observation model) produces noticeably different estimates of r_t and P_t compared to the modified approach using an extra-binomial observation model (Sect D of S1 Text).

Finally, both the basic and extra-binomial models perform poorly when fit to simulated weighted data, with only the weighted model being able to recover the true prevalence (Fig 3, column C). Unlike in the other models, estimates of growth rates are less prone to weighting-associated bias (Fig J in S1 Text). We emphasise caution when translating these results to real-world datasets: the simulated weighted model represents an extreme scenario where weights

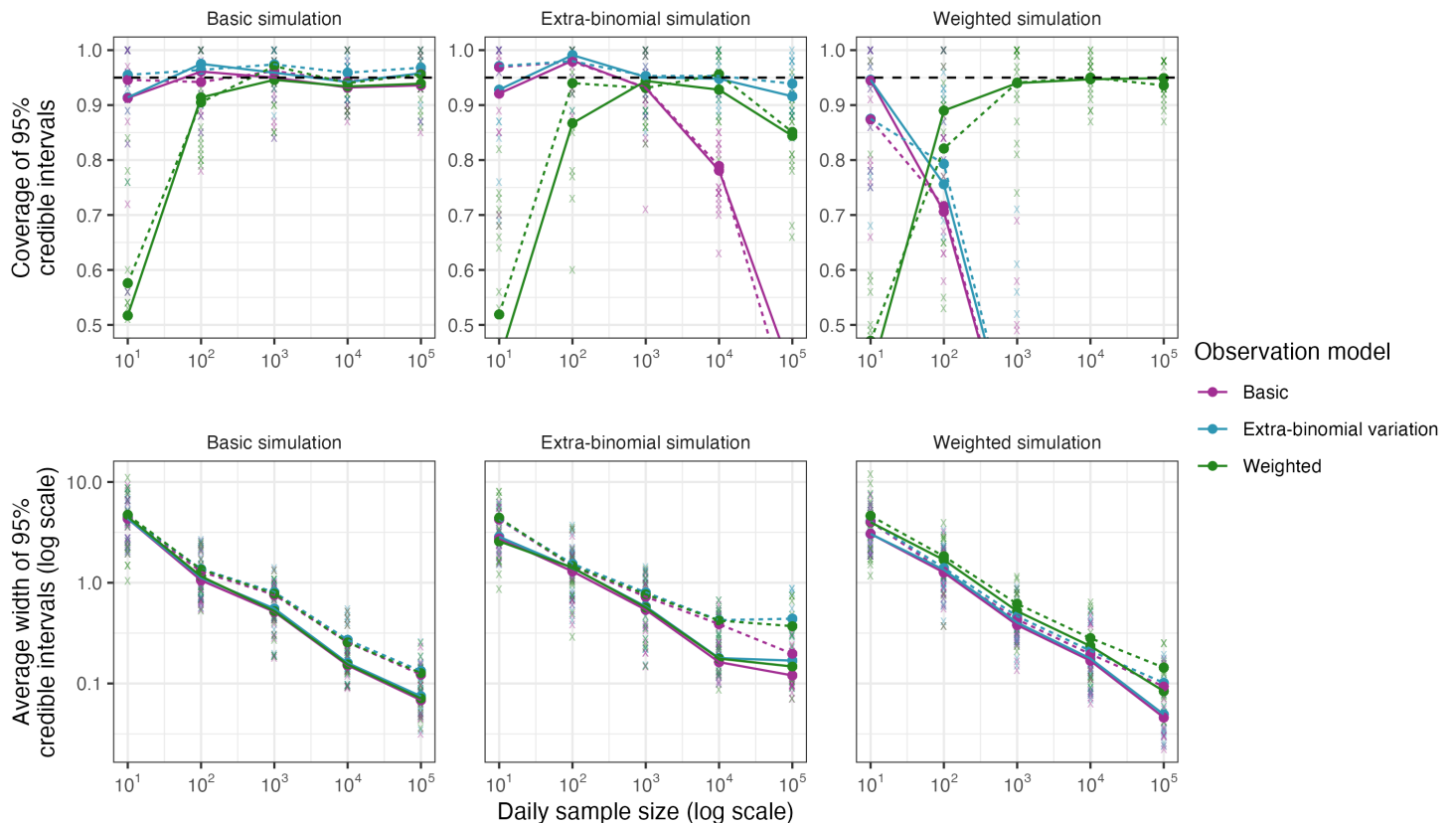


Fig 3. Coverage and average width of 95% credible intervals for prevalence P_t (%) from fitting all three observation models (purple: basic, blue: extra-binomial, green: weighted) to simulated data from each model (column A: basic, column B: extra-binomial, column C: weighted). Results from individual simulations are shown as semi-transparent crosses, with averages over 10 simulations shown as points connected by solid lines (for assumed $\sigma = 0.008$) and dashed lines (for assumed $\sigma = 0.016$). A range of assumed daily sample sizes n_t are considered (x-axis). The horizontal black dashed line indicates the target coverage of 0.95. The y-axis for coverage is truncated to (0.5,1.0), although the coverage in some cases falls outside this range: reaching a minimum average of 0.33 for the basic model fit to the extra-binomial simulations and a minimum average of 0 for the basic and extra-binomial models fit to the weighted simulations (all when $n_t = 10^5$).

<https://doi.org/10.1371/journal.pcbi.1013515.g003>

and individual swab positivity are perfectly correlated and the weights are known. In practice, these weights must be estimated through techniques such as random iterative method (RIM) weighting [40], which introduces additional uncertainty that is not accounted for by this model.

Comparing approaches

In this section, we compare the SIMPLE, Eales, and Abbott approaches on 10 simulated datasets from each model with default parameterisations (Table 1), assuming a beta-binomial observation process. All three models produce very similar posterior predictive distributions for swab positivity, and similar posterior distributions for prevalence, suggesting that, given the same observation model, the estimation of P_t and n_t^+ is robust to the differences in smoothing assumptions between models (Fig 4).

The models in the Eales and Abbott approaches both enforce a fixed minimum amount of smoothness in the data. In the case of the Eales approach, this is via spline knots being placed less frequently than the observed data; while in the case of the Abbott approach, this results from modelling prevalence as a convolution of smooth incidence and the test-sensitivity

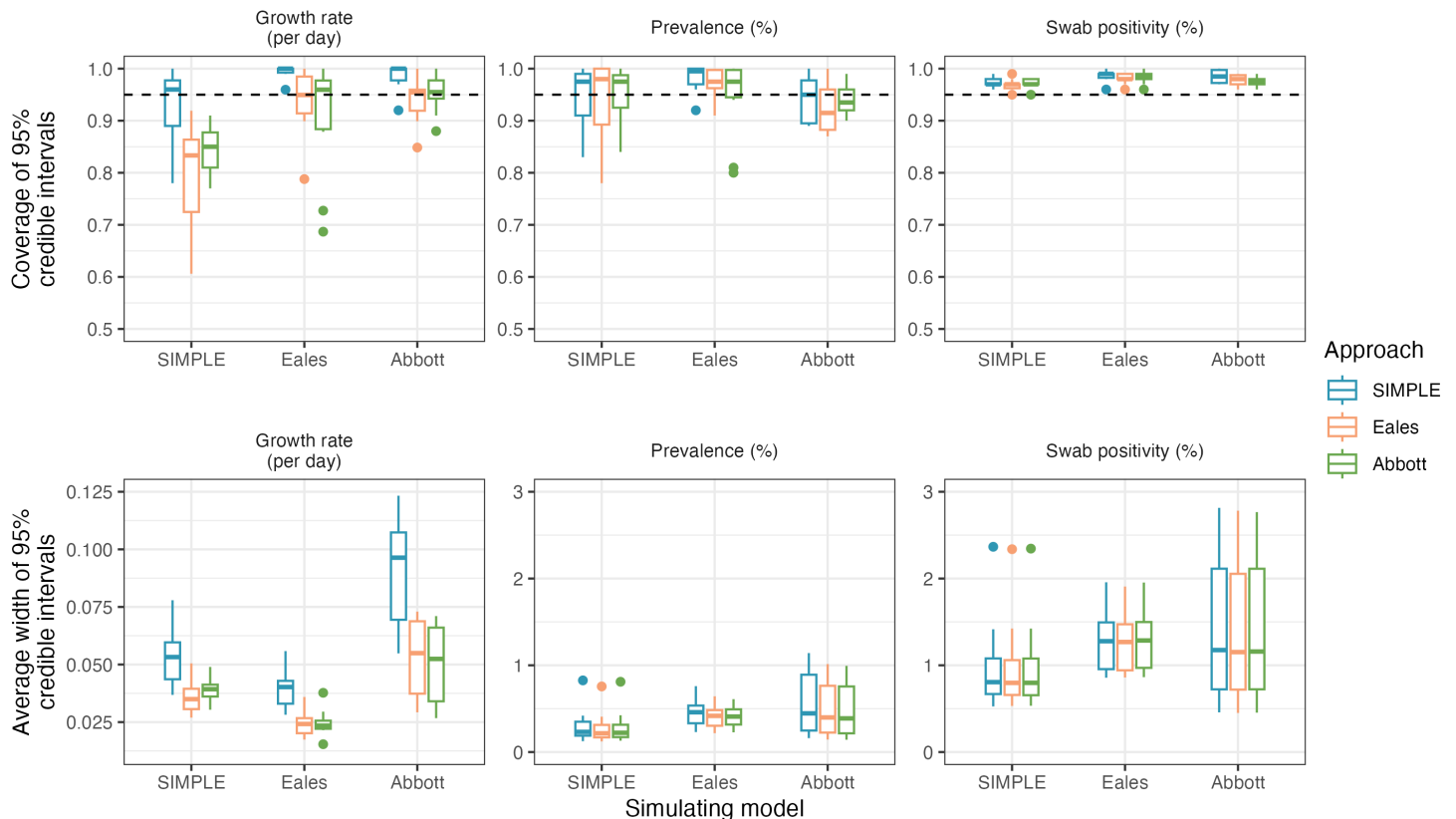


Fig 4. Coverage and average width of 95% credible intervals for the growth rate and prevalence, and coverage and average width of 95% predictive credible intervals for swab positivity. Each approach (SIMPLE, Eales, and Abbott) is fit to 10 simulated datasets from each model. The horizontal black dashed line indicates the target coverage of 0.95. Boxes present the interquartile range of the results with the median shown as a horizontal line. Whiskers extend to the most extreme data point within 1.5 times the interquartile range from the box. Outliers are shown as points.

<https://doi.org/10.1371/journal.pcbi.1013515.g004>

function and assuming that the incidence function is infinitely differentiable. This minimum smoothness results in the Eales and Abbott approaches producing narrower credible intervals on r_t than the SIMPLE approach, helpful when the true growth rate is smooth (the SIMPLE approach overcovers r_t in simulations from the Eales and Abbott models), but results in undercoverage when the true growth rate is more variable (Fig 4). These are further examples of simple modelling assumptions leading to “model-heavy” inferences.

The total time taken to fit the models to all 30 simulations was 18m 40s for the SIMPLE (extra-binomial) approach, 54m 58s for the Abbott approach, and 47m 44s for the Eales approach. A successful iteration of the Abbott approach takes less time than the Eales approach, however, the Abbott approach sometimes fails to converge, requiring refitting of the model. Further refinement of the code and/or better selection of prior distributions may improve convergence times for the Abbott approach. A more comprehensive comparison of the runtime of each approach is provided in Sect H of S1 Text.

The REACT-1 study

In this section we compare all three approaches on the REACT-1 dataset. First we focus on estimating the growth rate r_t and prevalence P_t , and then consider estimation of the reproduction number R_t . Assuming fixed values of the static parameters over the entire study period

of 700 days may not be appropriate [41]. To assess temporal variation in these parameters, we fit the models separately to four time periods: 1 May 2020 to 3 December 2020 (REACT-1 study rounds 1-to-7), 30 December 2020 to 12 July 2021 (REACT-1 study rounds 8-to-13), 9 September 2021 to 17 December 2021 (REACT-1 study rounds 14-to-16), and 5 January 2022 to 31 March 2022 (REACT-1 study rounds 17-to-19). These periods align approximately with changes in the dominant SARS-CoV-2 variant in England (Wildtype, Alpha, Delta, and Omicron) and large gaps in sampling.

All approaches (when using beta-binomial observation distributions) produce similar estimates of the growth rate r_t , prevalence P_t , and predictive swab positivity n_t^+/n_t , when fit to the REACT-1 dataset (Fig 5), with an average width of 95% credible intervals for P_t of 0.26 to 0.27 percentage points. Minor differences, largely arising from different prior assumptions about the initial growth rate r_0 , are observed in estimates at t near zero. All approaches also produce very similar estimates of the overdispersion parameter ρ , with posterior mean values of 1.9×10^{-4} to 2.0×10^{-4} (Table 2). Other parameters are not directly comparable. Finally, all

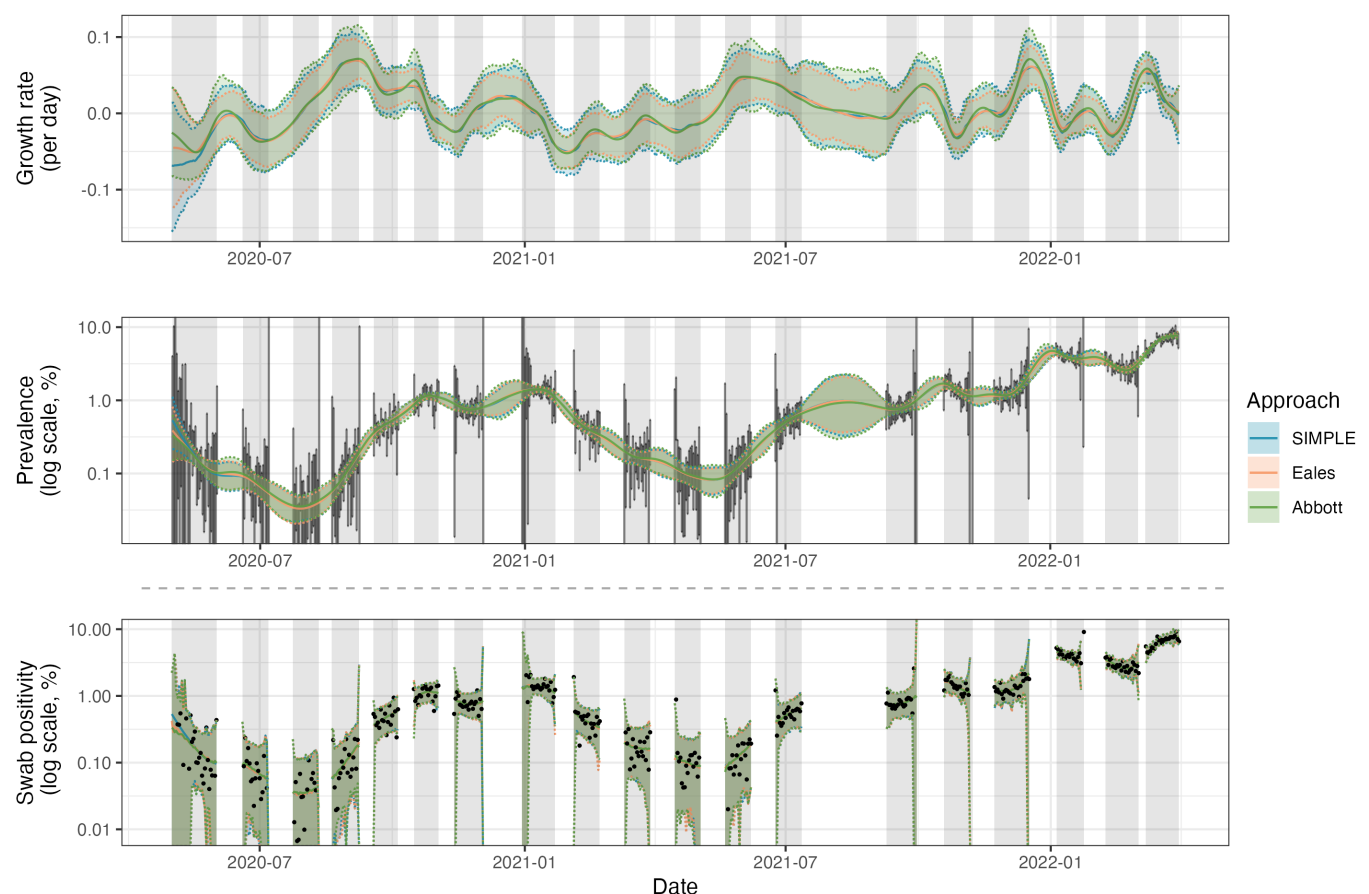


Fig 5. Estimates of the growth rate r_t , prevalence P_t , and predictive swab positivity n_t^+/n_t , for SARS-CoV-2 in England between 1 May 2020 and 31 March 2022 using data from the REACT-1 study. All three approaches are fit assuming a beta-binomial observation distribution. Solid coloured lines show the posterior means while shading and dashed lines show 95% credible intervals (of the posterior distribution for r_t and P_t , and of the posterior predictive distribution for n_t^+/n_t). Independent daily confidence intervals from the Agresti-Coull method [14] for P_t are shown in vertical grey lines. The data, daily observed swab positivity n_t^+/n_t , are shown in black points. Grey shading indicates the periods in which sampling was conducted. The predictive distribution for swab positivity depends on the number of swabs taken each day n_t , which tends to be lower in the early and late periods of each sampling round, hence the wider credible intervals at the boundaries of each study round.

<https://doi.org/10.1371/journal.pcbi.1013515.g005>

Table 2. Results from fitting the three approaches to data from the REACT-1 prevalence study. Runtimes were measured once for each dataset considered and can vary considerably. Convergence diagnostics of maximum $\hat{R}$ and minimum ESS are reported, although these are not directly comparable between the SIMPLE and Eales/Abbott approaches. Measures of fit are reported as coverage of the posterior predictive distribution and average width of 95% credible intervals on P_t (in terms of percentage points). Parameter estimates are shown as posterior means with 95% credible intervals in parentheses. Note that σ_{Eales} depends on the knot spacing, which varies slightly between study rounds, so these estimates are not directly comparable even within the same model.

Study round	All rounds	1-to-7	8-to-13	14-to-16	17-to-19
Observations	400	147	119	67	67
Duration (days)	700	217	195	100	86
SIMPLE approach					
Stopping criteria: max $\hat{R} < 1.05$ and min ESS > 100					
Runtime	4m 12s	42s	46s	25s	14s
Max $\hat{R}$ /Min ESS	1.01/113	1.02/110	1.02/100	1.04/110	1.04/129
Coverage of 95% Cr.I. n_t^+	97.2%	97.3%	97.5%	98.5%	98.5%
Avg. width of 95% Cr.I. P_t	0.27	0.16	0.17	0.37	0.84
Parameter σ	0.010 (0.0071, 0.011)	0.011 (0.0055, 0.022)	0.0083 (0.0041, 0.016)	0.015 (0.0068, 0.028)	0.015 (0.0075, 0.028)
Parameter ρ ($\times 10^{-4}$)	1.9 (1.3, 2.7)	2.6 (1.6, 4.5)	1.3 (0.61, 2.2)	1.1 (0.15, 3.1)	2.8 (1.2, 5.0)
SIMPLE approach (reproduction number epidemic model)					
Stopping criteria: max $\hat{R} < 1.05$ and min ESS > 100					
Runtime	23m 38s	3m 9s	1m 26s	59s	1m 23s
Max $\hat{R}$ /Min ESS	1.02/101	1.04/106	1.02/101	1.03/109	1.04/127
Coverage of 95% Cr.I. n_t^+	97.2%	96.6%	97.5%	98.5%	98.5%
Avg. width of 95% Cr.I. P_t	0.26	0.14	0.16	0.36	0.79
Parameter σ_R	0.069 (0.050, 0.097)	0.063 (0.034, 0.12)	0.056 (0.023, 0.11)	0.12 (0.046, 0.23)	0.11 (0.050, 0.22)
Parameter ρ ($\times 10^{-4}$)	2.0 (1.4, 2.6)	2.6 (1.8, 3.8)	1.3 (0.55, 2.2)	1.2 (0.12, 2.8)	3.4 (1.6, 6.0)
Eales approach*					
Stopping criteria: 500 samples (200 warm-up and 300 sampling)					
Runtime	38m 54s*	2m 58s	1m 14s	43s	35s
Max $\hat{R}$ /Min ESS	1.04/152*	1.02/380	1.02/259	1.01/245	1.01/260
Coverage of 95% Cr.I. n_t^+	96.5%	96.6%	97.5%	97.0%	98.5%
Avg. width of 95% Cr.I. P_t	0.26	0.14	0.15	0.34	0.74
Parameter σ_{Eales}	0.11 (0.084, 0.16)	0.097 (0.051, 0.17)	0.074 (0.033, 0.15)	0.15 (0.060, 0.30)	0.15 (0.079, 0.28)
Parameter ρ ($\times 10^{-4}$)	1.9 (1.4, 2.5)	2.7 (1.7, 4.1)	1.3 (0.59, 2.3)	0.98 (0.075, 2.6)	2.9 (1.2, 5.4)
Abbott approach					
Stopping criteria: 500 samples (200 warm-up and 300 sampling)					
Runtime	6m 43s	47s	49s	40s	29s
Max $\hat{R}$ /Min ESS	1.03/223	1.02/163	1.02/124	1.02/144	1.02/164
Coverage of 95% Cr.I. n_t^+	97.5%	98.0%	97.5%	95.5%	97.0%
Avg. width of 95% Cr.I. P_t	0.27	0.15	0.17	0.33	0.75
Parameter σ_{Abbott}	0.074 (0.033, 0.21)	0.069 (0.026, 0.19)	0.077 (0.024, 0.22)	0.10 (0.031, 0.23)	0.081 (0.026, 0.22)
Parameter ℓ	8.9 (0.38, 22)	22 (0.58, 79)	19 (0.29, 100)	3.9 (0.18, 20)	7.3 (0.17, 27)
Parameter ρ ($\times 10^{-4}$)	1.9 (1.3, 2.5)	2.7 (1.7, 4.1)	1.3 (0.54, 2.3)	1.0 (0.16, 2.5)	3.1 (1.5, 5.8)

*In order to obtain convergence of the Eales approach on the full REACT-1 dataset, the "maximum tree-depth" HMC hyperparameter was increased from 15 to 16, resulting in an increase in runtime.

<https://doi.org/10.1371/journal.pcbi.1013515.t002>

approaches produced posterior predictive distributions for n_t^+ that slightly overcovered the observed data (Table 2). These estimates differ from estimates reported by the REACT-1 study due to our use of a beta-binomial observation model, rather than the binomial observation model used in the original study [17] - we compare these for the Eales approach in Sect D of S1 Text.

Results from the SIMPLE approach (with the growth rate epidemic model) suggest that σ was higher in study rounds 14-to-16 and 17-to-19 (central estimates of $\sigma = 0.015$) than in study rounds 1-to-7 (central estimate of $\sigma = 0.011$) and rounds 8-to-13 (central estimate of $\sigma = 0.0083$), indicating greater variability in the growth rate in later study rounds. The results also suggest that ρ was higher in study rounds 1-to-7 and 17-to-19 (central estimates of 2.6 and 2.8, respectively) than in study rounds 8-to-13 and 14-to-16 (central estimates of 1.3 and 1.1), suggesting greater observation noise at the study's start and end. These trends are consistent with results from the SIMPLE approach with the reproduction number epidemic model, and the Eales and Abbott approaches (Table 2). While differences in parameter estimates are not statistically significant, estimates of the growth rates r_t do show some sensitivity to whether the models are fit separately to each time period or all together (Sect I of S1 Text).

The SIMPLE approach (with the growth rate epidemic model) and Abbott approach exhibit comparable runtimes, taking 4m 12s and 6m 43s on the full dataset (all study rounds), or 14s and 29s on the smallest dataset (study rounds 17-to-19), respectively. This is verified in Sect H of S1 Text, where a more extensive comparison of runtimes is provided. The Abbott approach can sometimes fail to converge, requiring refitting of the model, thus increasing the total time taken to fit the model. The Eales approach is slower than both of these approaches, particularly on larger datasets, taking 38m 54s to fit to the full dataset and 35s to fit to the smallest dataset. This slowdown is partially due to the need to increase the maximum tree-depth HMC hyperparameter from 15 to 16 for convergence on larger datasets. Finally, the SIMPLE approach with the reproduction number epidemic model is the slowest on smaller datasets, but remains faster than the Eales approach on larger ones (taking 23m 8s on the full dataset and 1m 23s on the smallest dataset). The runtimes reported in Table 2 are for a single run of each approach. A more extensive analysis of runtimes is included in Sect H of S1 Text.

Despite each approach producing similar estimates of r_t and P_t and a similar predictive distribution for n_t^+ (Fig 5), a number of arbitrary decisions were made when fitting these models. For example, we present estimates of r_t and P_t from fitting to all 19 study rounds simultaneously with constant static parameters, rather than grouping the data into shorter time periods. We also assume the extra-binomial observation model is appropriate, and we do not consider survey weights (as daily-applicable weights were not available). We test these specific modelling decisions in Sect I of S1 Text. Furthermore, we find evidence for variant-specific parameter values in Sect J of S1 Text. Finally, we compare estimates from all approaches with official consensus estimates of the growth rate of COVID-19 in England produced by the UK Health Security Agency (UKHSA) [42] in Sect K of S1 Text.

The reproduction number. While all approaches produce very similar estimates of P_t , and similar estimates of r_t when fit to the REACT-1 data, estimates of R_t differ more substantially (Fig 6). Additional assumptions are required when estimating R_t and this is a quantity known to be sensitive to such assumptions [43]. All models in this section are fit using the extra-binomial observation model.

The SIMPLE and Abbott approaches produce similar central estimates of R_t , with the Abbott approach producing wider credible intervals on average, reflecting the incorporation of uncertainty in the test-sensitivity function (where the SIMPLE approach uses the mean of the test-sensitivity function). Adapting the SIMPLE approach to allow for uncertainty in this

function is possible: by treating the test-sensitivity function as a parameter within PMMH, individual PCR positivity curves can be sampled from estimates in the literature and incorporated into the model. By storing accepted samples of the test-sensitivity function, it is possible to integrate out uncertainty about this function alongside other parameters. We leave the implementation of this to future work.

The Eales approach produces estimates of R_t that are delayed and oversmoothed in comparison to the SIMPLE and Abbott approaches (most apparent in Fig 6B), a direct result of the trailing-window approach to estimating R_t . The delay can be partially accounted for by shifting estimates by $\tau/2$ days [30] (Fig 6C), but the oversmoothing is inherent to the approach. Furthermore, as a trailing window of length $\tau = 14$ days does not include the entire generation time distribution or test-sensitivity duration, additional biases are introduced. Note that, while the Eales approach was used to estimate R_t in the published results from the REACT-1 study, our estimates may not align with these due to our use of a beta-binomial observation model instead of the original binomial model.

Despite these differences, there is no clear best approach. Most notably, estimates from the SIMPLE and Abbott approaches depend on the assumed test-sensitivity function, which

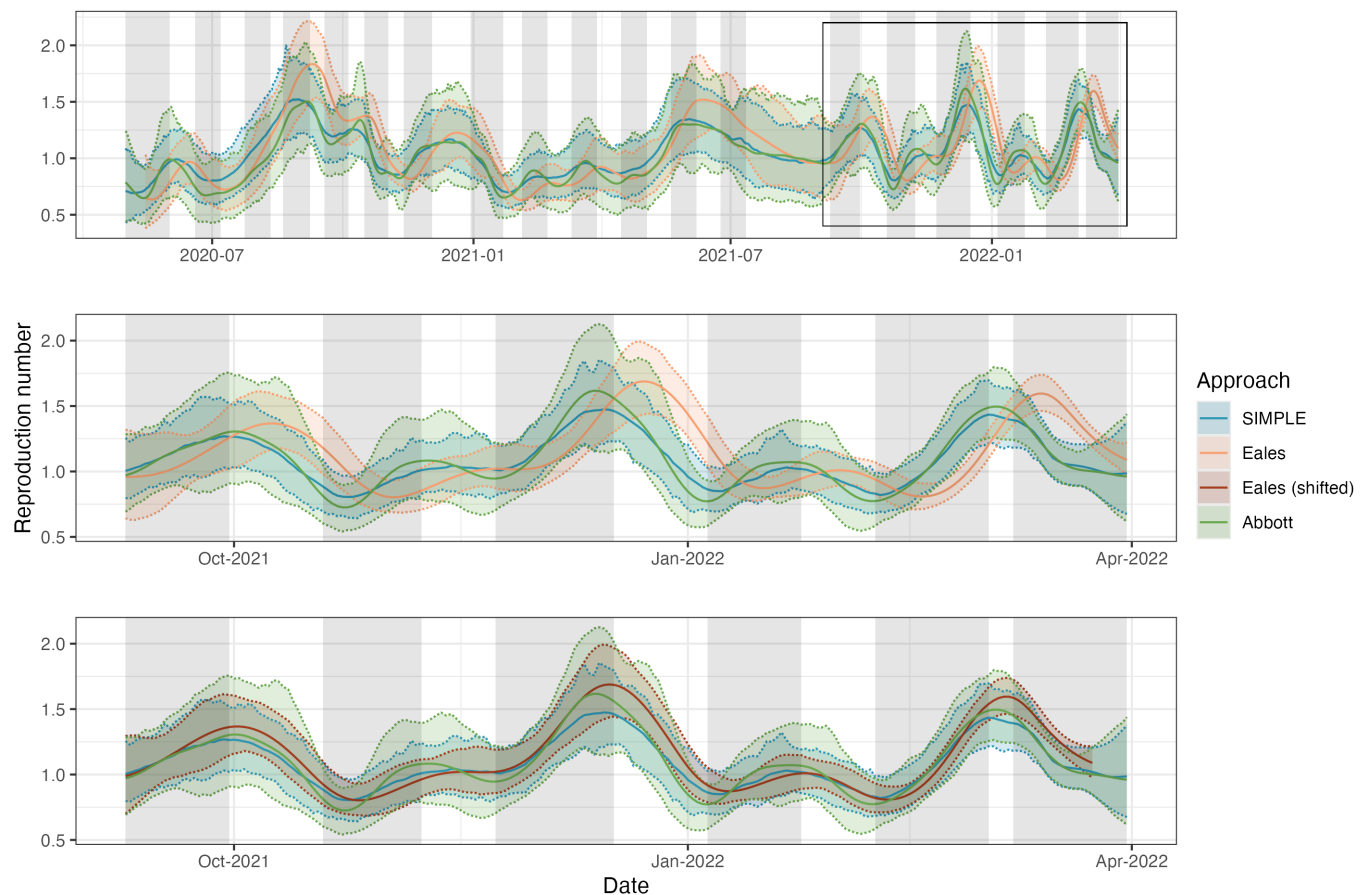


Fig 6. Estimates of the instantaneous reproduction number R_t for SARS-CoV-2 in England between 1 May 2020 and 31 March 2022 using data from the REACT-1 study. All approaches are fit assuming a beta-binomial observation distribution. Solid coloured lines show central estimates while shading and dashed lines show 95% credible intervals. Grey shading indicates the periods in which sampling was conducted. The second panel shows the same estimates for a shorter period (9 September 2021 to 31 March 2022), emphasising the differences between the three approaches. The third panel shows the same estimates, with estimates from the Eales approach shifted by $\tau/2 = 7$ days to partially account for bias induced by the trailing-window smoothing assumption.

<https://doi.org/10.1371/journal.pcbi.1013515.g006>

can vary significantly between settings. This is explored further in Sect E of [S1 Text](#), where we compare three different curves estimated for SARS-CoV-2 RT-PCR tests, and demonstrate how the choice of curve impacts estimates of R_t . We also compare estimates of R_t from these approaches with official consensus estimates produced by the UKHSA [42] in Sect K of [S1 Text](#).

As the reproduction number epidemic model in the SIMPLE approach is no longer Markovian, we must use fixed-lag resampling during the parameter inference stage [16], leading to an increase in computation time (23m 38s for the full dataset or 1m 23s on the smallest dataset). The faster runtime of the SIMPLE approach with the growth rate epidemic model could be obtained by fitting a Markovian model (instead of the renewal model), such as a compartmental susceptible-infectious-recovered-type model, which has the added benefit of not requiring a test-sensitivity function. However, this type of model places additional assumptions on the underlying epidemic dynamics, which can be difficult to validate. We leave this for future work.

Discussion

Smoothing prevalence data allows multiple days to inform point estimates, reducing noise and increasing the confidence in estimates. However, all smoothing methods make assumptions, whether explicitly or implicitly, about the underlying process, and results can be sensitive to these assumptions.

A key concern highlighted in this paper is the presence of overdispersion in the observed data. As shown in Sect 3.1, not properly accounting for this can lead to biased prevalence estimates and credible intervals with poor coverage. Growth rate estimates also become more variable, as the model attempts to explain extra-binomial noise through artificial variability in r_t . In addition to the simulated results presented in this paper, we also find that this impacts real-world estimates of P_t and r_t in the REACT-1 study. An explicit comparison is given in Sect D of [S1 Text](#), emphasising that results can vary considerably, particularly in earlier study rounds. Once overdispersion is appropriately handled, all approaches produce accurate prevalence estimates, even when fit to data generated by other models. This is because prevalence estimates are largely driven by the observed data, not by the smoothing assumptions.

Even when the data are not generated by a beta-binomial sampling process, the overdispersion parameter ρ serves as a general error term that can capture unmodelled heterogeneities and sources of variance. The similarity of estimates of ρ between the approaches (when fit to real-world data from the REACT-1 study) suggests that this parameter is estimable from observed data independently of the assumed smoothing mechanism. Growth rate estimates (and to a lesser extent, prevalence estimates) in the REACT-1 study were sensitive to the inclusion of this parameter. As the allowance for overdispersion, even when simulated data were generated from a binomial model, had little impact on the accuracy of estimates, we recommend including this parameter in all models. If survey weights are available and the daily sample size is sufficiently large (typically $n_t > 100$, although low-prevalence scenarios may require more), then the weighted model is able to account for both overdispersion and sampling bias, and thus should be the first choice. Unfortunately, without access to survey weights applicable on a daily basis, we were unable to apply this model to the REACT-1 study.

In addition to smoothing prevalence data, the approaches presented in this paper allow for the estimation of the growth rate in prevalence r_t . Unlike estimates of prevalence, assumed smoothing dynamics have a substantial impact on estimates of r_t . Only the SIMPLE approach produced 95% credible intervals with consistently good coverage of this quantity, at the cost

of substantially wider intervals. However, as prevalence is naturally a convolution of past incidence, the additional smoothness implied by the SIMPLE approach with the reproduction number epidemic model, or the Abbott approach, may provide an appropriate way to reduce estimated uncertainty about r_t .

Given a predetermined model structure, uncertainty about r_t and P_t arises from two sources: (1) process and observation noise in the epidemic and observation model, and (2) uncertainty about the static parameters governing these models. While uncertainty associated with the former can only practically be reduced by collecting more data, uncertainty about static parameters can be decreased by using more informative prior distributions. At smaller values of n_t , a wide uniform prior distribution can result in the posterior distribution assigning probability mass to implausibly large values of σ and ρ , leading to overestimation of the width of credible intervals for r_t and P_t . Using more informative prior distributions that place less prior mass on implausibly large values can help to reduce the width of the credible intervals for r_t , P_t , and n_t^+ , particularly at smaller values of n_t . This is demonstrated in Sect A of [S1 Text](#).

Estimating R_t from prevalence data requires assumptions about the relationship between incidence and swab positivity. This can be achieved by assuming R_t is constant over a sufficiently long time period (Eales approach), by incorporating a test-sensitivity function (SIMPLE and Abbott approaches), or by fitting a fully mechanistic model (not considered here). Each approach introduces different sources of potential bias. The trailing-window assumption in the Eales approach guarantees some degree of misspecification [29], which can only be partially mitigated by shifting estimates by $\tau/2$ days [30]. Test-sensitivity functions offer an alternative but also vary substantially: several distinct estimates of this function for SARS-CoV-2 RT-PCR tests were produced during the COVID-19 pandemic [24,37,44,45], with meaningful differences in both duration and peak sensitivity. Ideally this function would be estimated for each epidemiological setting being modelled, which would require repeat swabs from a subset of survey participants, although collecting sufficient data may be infeasible in low-prevalence scenarios.

An alternative way to relate incidence to prevalence is to model individual-level cycle threshold (Ct) values [38]. Instead of relying on an assumed test-sensitivity function, this approach fits a model of Ct values as a function of time since infection, replacing the observation model for aggregated data with one for individual Ct data. Models for individual Ct value trajectories have previously been estimated from cross-sectional data, without needing repeated sampling from individuals. This method makes fuller use of the information contained in Ct values, rather than reducing test results to binary outcomes, but requires access to individual-level data, is computationally intensive, and depends on correctly specifying the Ct model.

Our SIMPLE approach requires no external software (e.g. Stan), has a faster runtime than the Eales approach, and avoids the Gaussian process approximations used in the Abbott approach. This computational efficiency allows for the rapid testing of a range of models, as demonstrated in Sect I of [S1 Text](#), where we use the SIMPLE approach to fit a range of models to the REACT-1 dataset. We also demonstrate an adaptation that allows the SIMPLE approach to estimate variant-specific growth rates while leveraging all collected data (including unsequenced samples) in Sect J of [S1 Text](#). Finally, the sequential nature of SMC allows for easy modification for online inference, where replacing PMMH with an SMC² algorithm [46] would allow for real-time inference as new data are collected without the need to refit the model from scratch.

There are two main limitations to the SIMPLE approach. First, the resampling step in the PMMH algorithm is computationally expensive. For Markovian models, past-state resampling can be disabled during PMMH, enabling fast inference. For non-Markovian models, past-state resampling is necessary, slowing down the PMMH algorithm, as seen in the SIMPLE model for the reproduction number. In these cases, the Abbott approach has a faster run-time. Second, PMMH relies on stochastic likelihood estimates, whose variance increases with data length and model complexity. This results in $O(T^2)$ time complexity for parameter inference, which was not a bottleneck for the datasets used here ($T \leq 700$) but will become important for longer time series. While we use relatively simple models in this paper, more complex models would benefit from more advanced algorithms [47].

While we make several improvements to the Eales and Abbott approaches - including using beta-binomial observation distributions and modifying HMC hyperparameters - further optimisation is likely possible which could improve their relative performance. In particular, the Abbott approach sometimes fails to converge. Reparameterising the model, or using better-specified initial values, could reduce or prevent the occurrence of this. In the Abbott approach, we also recommend considering alternative, less smooth, Gaussian process kernels, as the squared exponential kernel is known to produce overly smooth estimates [48]. The very smooth kernel used in the Abbott approach can be partially compensated for by estimating a more variable growth rate r_t , which is why simulated data from the Abbott model (Fig 1) exhibits r_t values of greater magnitude than the SIMPLE and Eales approaches, on average.

Estimate precision can be improved by leveraging additional data sources. The original Abbott approach, for example, included an observation model for antibody data [7]. If individual-level data are available, the approach of Pouwels et al. [9], later refined in [49], can be used to model individual probabilities of testing positive. This approach performs post-stratification of individual-level estimates by demographic-geographical response types and has several advantages: multilevel regression and poststratification (MRP) has been shown to outperform classical survey weighting [50], and partial pooling improves estimation of demographic effects (whereas the approaches presented in this paper must produce demographic-specific estimates by fitting separate models). Future work could explore how to combine the approaches and lessons presented in this paper with MRP-type approaches to improve estimates of P_t and r_t .

Our findings indicate that while all three approaches provide robust estimates of prevalence and observed positive swabs, minor differences in their smoothing assumptions can impact growth rate estimates. Applying these approaches to data from the REACT-1 study, we observed that all three produced similar estimates of prevalence and growth rates, with the SIMPLE and Abbott approaches demonstrating greater computational efficiency. By presenting these approaches in common and general notation, alongside well-documented code online, this paper offers a suite of validated tools that can be readily adapted for future epidemic surveys.

Supporting information

S1 Text. Supplementary text. Supplementary methods, results, tables, and figures. (PDF)

Author contributions

Conceptualization: Nicholas Steyn, Marc Chadeau-Hyam, Christl A. Donnelly.

Data curation: Nicholas Steyn, Marc Chadeau-Hyam, Paul Elliott, Christl A. Donnelly.

Formal analysis: Nicholas Steyn, Marc Chadeau-Hyam, Christl A. Donnelly.

Funding acquisition: Marc Chadeau-Hyam, Paul Elliott, Christl A. Donnelly.

Investigation: Nicholas Steyn, Marc Chadeau-Hyam, Christl A. Donnelly.

Methodology: Nicholas Steyn, Marc Chadeau-Hyam, Christl A. Donnelly.

Project administration: Marc Chadeau-Hyam, Paul Elliott, Christl A. Donnelly.

Resources: Nicholas Steyn, Marc Chadeau-Hyam, Christl A. Donnelly.

Software: Nicholas Steyn.

Supervision: Marc Chadeau-Hyam, Christl A. Donnelly.

Validation: Nicholas Steyn, Marc Chadeau-Hyam, Christl A. Donnelly.

Visualization: Nicholas Steyn, Marc Chadeau-Hyam, Christl A. Donnelly.

Writing – original draft: Nicholas Steyn, Marc Chadeau-Hyam, Christl A. Donnelly.

Writing – review & editing: Nicholas Steyn, Marc Chadeau-Hyam, Paul Elliott, Christl A. Donnelly.

References

1. Murray J, Cohen AL. Infectious disease surveillance. *International Encyclopedia of Public Health*. Elsevier; 2017. p. 222–9. <https://doi.org/10.1016/b978-0-12-803678-5.00517-8>
2. Dutta H, Kaushik G, Dutta V. Wastewater-based epidemiology: a new frontier for tracking environmental persistence and community transmission of COVID-19. *Environ Sci Pollut Res Int*. 2022;29(57):85688–99. <https://doi.org/10.1007/s11356-021-17419-0> PMID: 34762243
3. Bengtsson L, Gaudart J, Lu X, Moore S, Wetter E, Sallah K, et al. Using mobile phone data to predict the spatial spread of cholera. *Sci Rep*. 2015;5:8923. <https://doi.org/10.1038/srep08923> PMID: 25747871
4. COVID-19 Community Mobility Report. <https://www.google.com/covid19/mobility?hl=en>
5. Broniatowski DA, Paul MJ, Dredze M. National and local influenza surveillance through Twitter: an analysis of the 2012–2013 influenza epidemic. *PLoS One*. 2013;8(12):e83672. <https://doi.org/10.1371/journal.pone.0083672> PMID: 24349542
6. Bradley VC, Kuriwaki S, Isakov M, Sejdinovic D, Meng X-L, Flaxman S. Unrepresentative big surveys significantly overestimated US vaccine uptake. *Nature*. 2021;600(7890):695–700. <https://doi.org/10.1038/s41586-021-04198-4> PMID: 34880504
7. Abbott S, Funk S. Estimating epidemiological quantities from repeated cross-sectional prevalence measurements. Cold Spring Harbor Laboratory. 2022. <https://doi.org/10.1101/2022.03.29.22273101> <http://dx.doi.org/10.1101/2022.03.29.22273101>
8. Elliott P, Whitaker M, Tang D, Eales O, Steyn N, Bodinier B, et al. Design and Implementation of a National SARS-CoV-2 monitoring program in England: REACT-1 Study. *Am J Public Health*. 2023;113(5):545–54. <https://doi.org/10.2105/AJPH.2023.307230> PMID: 36893367
9. Pouwels KB, House T, Pritchard E, Robotham JV, Birrell PJ, Gelman A, et al. Community prevalence of SARS-CoV-2 in England from April to November, 2020: results from the ONS coronavirus infection survey. *Lancet Public Health*. 2021;6(1):e30–8. [https://doi.org/10.1016/S2468-2667\(20\)30282-6](https://doi.org/10.1016/S2468-2667(20)30282-6) PMID: 33308423
10. Ward T, Fyles M, Glaser A, Paton RS, Ferguson W, Overton CE. The real-time infection hospitalisation and fatality risk across the COVID-19 pandemic in England. *Nat Commun*. 2024;15(1):4633. <https://doi.org/10.1038/s41467-024-47199-3> PMID: 38821930
11. Eales O, Haw D, Wang H, Atchison C, Ashby D, Cooke GS, et al. Dynamics of SARS-CoV-2 infection hospitalisation and infection fatality ratios over 23 months in England. *PLoS Biol*. 2023;21(5):e3002118. <https://doi.org/10.1371/journal.pbio.3002118> PMID: 37228015
12. Chadeau-Hyam M, Wang H, Eales O, Haw D, Bodinier B, Whitaker M, et al. SARS-CoV-2 infection and vaccine effectiveness in England (REACT-1): a series of cross-sectional random community surveys. *Lancet Respir Med*. 2022;10(4):355–66. [https://doi.org/10.1016/S2213-2600\(21\)00542-7](https://doi.org/10.1016/S2213-2600(21)00542-7) PMID: 35085490

13. Total cost of the COVID-19 infection survey. Office for National Statistics. <https://www.ons.gov.uk/aboutus/transparencyandgovernance/freedomofinformationfoi/totalcostofthecovid19infectionsurvey>
14. Agresti A, Coull BA. Approximate is better than “exact” for interval estimation of binomial proportions. *The American Statistician*. 1998;52(2):119. <https://doi.org/10.2307/2685469>
15. Doucet A, De Freitas N, Gordon N. *Sequential Monte Carlo Methods in Practice*. Doucet A, De Freitas N, Gordon N, editors. New York: Springer. 2001.
16. Steyn N, Parag KV, Thompson RN, Donnelly CA. A primer on inference and prediction with epidemic renewal models and sequential Monte Carlo. *Stat Med*. 2025;44(18–19):e70204. <https://doi.org/10.1002/sim.70204> PMID: 40772730
17. Eales O, Ainslie KEC, Walters CE, Wang H, Atchison C, Ashby D, et al. Appropriately smoothing prevalence data to inform estimates of growth rate and reproduction number. *Epidemics*. 2022;40:100604. <https://doi.org/10.1016/j.epidem.2022.100604> PMID: 35780515
18. Creswell R, Robinson M, Gavaghan D, Parag KV, Lei CL, Lambert B. A Bayesian nonparametric method for detecting rapid changes in disease transmission. *J Theor Biol*. 2023;558:111351. <https://doi.org/10.1016/j.jtbi.2022.111351> PMID: 36379231
19. Parag KV, Thompson RN, Donnelly CA. Are epidemic growth rates more informative than reproduction numbers?. *J R Stat Soc Ser A Stat Soc*. 2022;:10.1111/rssa.12867. <https://doi.org/10.1111/rssa.12867> PMID: 35942192
20. Cori A, Ferguson NM, Fraser C, Cauchemez S. A new framework and software to estimate time-varying reproduction numbers during epidemics. *Am J Epidemiol*. 2013;178(9):1505–12. <https://doi.org/10.1093/aje/kwt133> PMID: 24043437
21. Charniga K, Park SW, Akhmetzhanov AR, Cori A, Dushoff J, Funk S, et al. Best practices for estimating and reporting epidemiological delay distributions of infectious diseases. *PLoS Comput Biol*. 2024;20(10):e1012520. <https://doi.org/10.1371/journal.pcbi.1012520> PMID: 39466727
22. Ogi-Gittins I, Hart WS, Song J, Nash RK, Polonsky J, Cori A, et al. A simulation-based approach for estimating the time-dependent reproduction number from temporally aggregated disease incidence time series data. *Epidemics*. 2024;47:100773. <https://doi.org/10.1016/j.epidem.2024.100773> PMID: 38781911
23. Xu X, Wu Y, Kummer AG, Zhao Y, Hu Z, Wang Y, et al. Assessing changes in incubation period, serial interval, and generation time of SARS-CoV-2 variants of concern: a systematic review and meta-analysis. *BMC Med*. 2023;21(1):374. <https://doi.org/10.1186/s12916-023-03070-8> PMID: 37775772
24. Hellewell J, Russell TW, SAFER Investigators and Field Study Team, Crick COVID-19 Consortium, CMMID COVID-19 Working Group, Beale R, et al. Estimating the effectiveness of routine asymptomatic PCR testing at different frequencies for the detection of SARS-CoV-2 infections. *BMC Med*. 2021;19(1):106. <https://doi.org/10.1186/s12916-021-01982-x> PMID: 33902581
25. Luceño A, Ceballos FD. Describing extra-binomial variation with partially correlated models. *Communications in Statistics - Theory and Methods*. 1995;24(6):1637–53. <https://doi.org/10.1080/03610929508831576>
26. Wagner B, Riggs P, Mikulich-Gilbertson S. The importance of distribution-choice in modeling substance use data: a comparison of negative binomial, beta binomial, and zero-inflated distributions. *Am J Drug Alcohol Abuse*. 2015;41(6):489–97. <https://doi.org/10.3109/00952990.2015.1056447> PMID: 26154448
27. Gelman A, Rubin DB. Inference from iterative simulation using multiple sequences. *Statist Sci*. 1992;7(4). <https://doi.org/10.1214/ss/1177011136>
28. Bezanson J, Edelman A, Karpinski S, Shah VB. Julia: a fresh approach to numerical computing. *SIAM Rev*. 2017;59(1):65–98. <https://doi.org/10.1137/141000671>
29. Steyn N, Parag KV. Robust uncertainty quantification in popular estimators of the instantaneous reproduction number. *Am J Epidemiol*. 2025;:kwaf165. <https://doi.org/10.1093/aje/kwaf165> PMID: 40754791
30. Gostic KM, McGough L, Baskerville EB, Abbott S, Joshi K, Tedijanto C, et al. Practical considerations for measuring the effective reproductive number, Rt. *PLoS Comput Biol*. 2020;16(12):e1008409. <https://doi.org/10.1371/journal.pcbi.1008409> PMID: 33301457
31. R Core Team. 2021.
32. Gabry J, Češnovar R, Johnson A, Bröder S. 2024.
33. Solin A, Särkkä S. Hilbert space methods for reduced-rank Gaussian process regression. *Stat Comput*. 2019;30(2):419–46. <https://doi.org/10.1007/s11222-019-09886-w>
34. Riutort-Mayol G, Bürkner P-C, Andersen MR, Solin A, Vehtari A. Practical Hilbert space approximate Bayesian Gaussian processes for probabilistic programming. *Stat Comput*. 2022;33(1). <https://doi.org/10.1007/s11222-022-10167-2>

35. Kanagawa M, Hennig P, Sejdinovic D, Sriperumbudur BK. Gaussian processes and Kernel methods: a review on connections and equivalences. 2018.
36. Calistri P, Amato L, Puglia I, Cito F, Di Giuseppe A, Danzetta ML, et al. Infection sustained by lineage B.1.1.7 of SARS-CoV-2 is characterised by longer persistence and higher viral RNA loads in nasopharyngeal swabs. *Int J Infect Dis.* 2021;105:753–5. <https://doi.org/10.1016/j.ijid.2021.03.005> PMID: 33684558
37. Binny RN, Priest P, French NP, Parry M, Lustig A, Hendy SC, et al. Sensitivity of reverse transcription polymerase chain reaction tests for severe acute respiratory syndrome coronavirus 2 through time. *J Infect Dis.* 2022;227(1):9–17. <https://doi.org/10.1093/infdis/jiac317> PMID: 35876500
38. Hay JA, Kennedy-Shaffer L, Kanjilal S, Lennon NJ, Gabriel SB, Lipsitch M, et al. Estimating epidemiologic dynamics from cross-sectional viral load distributions. *Science.* 2021;373(6552):eabh0635. <https://doi.org/10.1126/science.abh0635> PMID: 34083451
39. Harrell FEJ. Hmisc: Harrell miscellaneous. 2024.
40. Sharot T. Weighting survey results. *Journal of the Market Research Society.* 1986;28(3):269–84.
41. Cazelles B, Champagne C, Dureau J. Accounting for non-stationarity in epidemiology by embedding time-varying parameters in stochastic models. *PLoS Comput Biol.* 2018;14(8):e1006211. <https://doi.org/10.1371/journal.pcbi.1006211> PMID: 30110322
42. Agency UHS. The R value and growth rate. 2022. <https://www.gov.uk/guidance/the-r-value-and-growth-rate>
43. Brockhaus EK, Wolfram D, Stadler T, Osthege M, Mitra T, Littek JM, et al. Why are different estimates of the effective reproductive number so different? A case study on COVID-19 in Germany. *PLoS Comput Biol.* 2023;19(11):e1011653. <https://doi.org/10.1371/journal.pcbi.1011653> PMID: 38011276
44. Kucirka LM, Lauer SA, Laeyendecker O, Boon D, Lessler J. Variation in false-negative rate of reverse transcriptase polymerase chain reaction-based SARS-CoV-2 tests by time since exposure. *Ann Intern Med.* 2020;173(4):262–7. <https://doi.org/10.7326/M20-1495> PMID: 32422057
45. Zhang Z, Bi Q, Fang S, Wei L, Wang X, He J, et al. Insight into the practical performance of RT-PCR testing for SARS-CoV-2 using serological data: a cohort study. *Lancet Microbe.* 2021;2(2):e79–87. [https://doi.org/10.1016/S2666-5247\(20\)30200-7](https://doi.org/10.1016/S2666-5247(20)30200-7) PMID: 33495759
46. Chopin N, Jacob PE, Papaspiliopoulos O. SMC2: an efficient algorithm for sequential analysis of state space models. *Journal of the Royal Statistical Society Series B: Statistical Methodology.* 2012;75(3):397–426. <https://doi.org/10.1111/j.1467-9868.2012.01046.x>
47. Kantas N, Doucet A, Singh SS, Maciejowski J, Chopin N. On particle methods for parameter estimation in state-space models. *Statist Sci.* 2015;30(3). <https://doi.org/10.1214/14-sts511>
48. Hadji A, Szabó B. Can we trust Bayesian Uncertainty Quantification from Gaussian process priors with squared exponential covariance Kernel?. *SIAM/ASA J Uncertainty Quantification.* 2021;9(1):185–230. <https://doi.org/10.1137/19m1253010>
49. Pouwels KB, Eyre DW, House T, Aspey B, Matthews PC, Stoesser N, et al. Improving the representativeness of UK's national COVID-19 Infection Survey through spatio-temporal regression and post-stratification. *Nat Commun.* 2024;15(1):5340. <https://doi.org/10.1038/s41467-024-49201-4> PMID: 38914564
50. Gelman A, Lax J, Phillips J, Gabry J, Trangucci R. Using multilevel regression and poststratification to estimate dynamic public opinion. Columbia University; 2018.