



qPCR data analysis: Better results through iconoclasm

Joel Tellinghuisen^{a,*}, Andrej-Nikolai Spiess^b

^a Department of Chemistry, Vanderbilt University Nashville, TN, 37235, USA

^b Department of Andrology, University Hospital Hamburg-Eppendorf, Hamburg, Germany

ARTICLE INFO

Handled by Jim Huggett

Keywords:

qPCR

Data analysis

Weighted least squares

Statistical errors

Chi-square

Calibration

ABSTRACT

The standard approach for quantitative estimation of genetic materials with qPCR is calibration with known concentrations for the target substance, in which estimates of the quantification cycle (C_q) are fitted to a straight-line function of $\log(N_0)$, where N_0 is the initial number of target molecules. The location of C_q for the unknown on this line then yields its N_0 . The most widely used definition for C_q is an absolute threshold that falls in the early growth cycles. This usage is flawed as commonly implemented: threshold set very close to the baseline level, which is estimated separately, from designated "baseline cycles." The absolute threshold is especially poor for dealing with the scale variability often observed for growth profiles. Scale-independent markers, like the first derivative maximum (FDM) and a relative threshold (C_r) avoid this problem. We describe improved methods for estimating these and other C_q markers and their standard errors, from a nonlinear algorithm that fits growth profiles to a 4-parameter log-logistic function plus a baseline function. By examining six multidilution, multi-replicate qPCR data sets, we find that nonlinear expressions are often preferred statistically for the dependence of C_q on $\log(N_0)$. This means that the amplification efficiency E depends on N_0 , in violation of another tenet of qPCR analysis. Neglect of calibration nonlinearity leads to biased estimates of the unknown. By logic, E estimates from calibration fitting pertain to the earliest baseline cycles, *not* the early growth cycles used to estimate E from growth profiles for single reactions. This raises concern about the use of the latter in lengthy extrapolations to estimate N_0 . Finally, we observe that replicate ensemble standard deviations greatly exceed predictions, implying that much better results can be achieved from qPCR through better experimental procedures, which likely include reducing pipette volume uncertainty.

1. Introduction

The goal of quantitative polymerase chain reaction (qPCR) is, as advertised, the quantification of small amounts of targeted genetic material through amplification to easily detected quantities [1]. Quantification may be relative to a chosen reference substance [2,3] or absolute. The standard approach for the latter is through calibration procedures that compare the unknown with results for the same substance measured at a range of known concentrations appropriate for the unknown [4–6]. Fig. 1 illustrates typical qPCR growth profiles obtained at 5 concentrations spanning 4 orders of magnitude in template copy number. Calibration with such data, like all classical calibration

procedures, involves identifying some property that depends monotonically on the concentration, fitting measurements of that property to a suitable response function, $z = f(x)$, and then solving the equation,

$$z_0 = f(x_0), \quad (1)$$

for the unknown x_0 , where z_0 is its measured response. Although analysts prefer linear response functions, there is no requirement for such, and calculation of x_0 and its uncertainty is a simple computational task for any $f(x)$ [8,9]. In qPCR the chosen calibration relationship is C_q vs. $\log(N_0)$, as illustrated in Fig. 2 for the data in Fig. 1. This choice is based on the exponential growth equation that is assumed to hold throughout the baseline region and into the early growth phase,

Abbreviations: qPCR, quantitative polymerase chain reaction; y and y_0 , fluorescence signal above baseline at cycle x and at cycle 0; E , amplification efficiency; C_q , quantification cycle; y_q , signal at $x = C_q$; N_0 , initial number of target molecules in sample; C_t , threshold cycle, where $y = y_q$; FDM and SDM, cycles where y reaches its maximal first and second derivatives, respectively; $Cy0$, intersection of a straight line tangent to the curve at the FDM with the baseline-corrected x -axis; LS, least squares; χ^2 , chi-square; w_i , statistical weight for i th data point; σ^2 and σ , variance and standard deviation; S , sum of weighted, squared residuals (= "Chisq" in KaleidaGraph fit results, = χ^2 when $w_i = 1/\sigma_i^2$); ν , statistical degrees of freedom, = # of data points – # of adjustable parameters; SD, standard deviation; SE, parameter standard error

* Corresponding author.

E-mail address: joel.tellinghuisen@vanderbilt.edu (J. Tellinghuisen).

<https://doi.org/10.1016/j.bdq.2019.100084>

Received 2 July 2018; Received in revised form 18 January 2019; Accepted 25 February 2019

Available online 05 June 2019

2214-7535/ © 2019 The Authors. Published by Elsevier GmbH. This is an open access article under the CC BY-NC-ND license (<http://creativecommons.org/licenses/by-nc-nd/4.0/>).

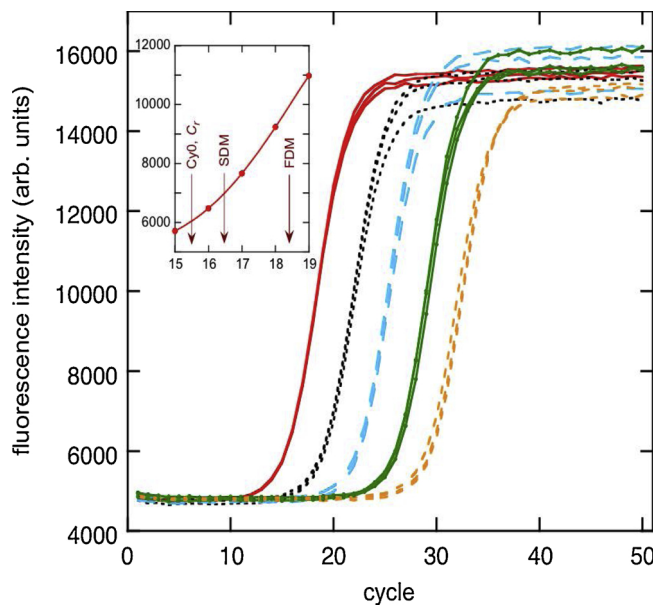


Fig. 1. qPCR fluorescence curves for lambda gDNA for 10-fold dilution from 188,000 copy numbers to 19, as recorded in triplicate by Rutledge and Stewart [7]. Inset shows positions of C_q markers for one reaction at highest concentration. With the threshold set at 12% of the (plateau – baseline) difference, the relative threshold C_r coincides with Cy0 within 0.1 cycle.

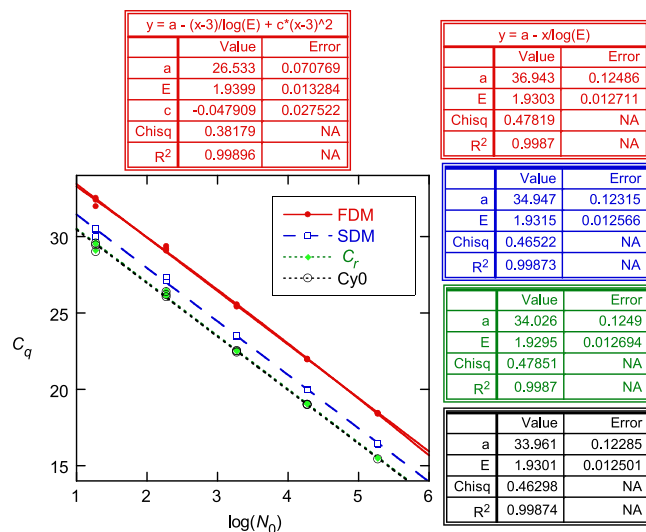


Fig. 2. Results of LS fits of 4 C_q markers from growth profiles in Fig. 1 to linear relation (right) and of FDM to quadratic centered at $\log(N_0) = 3$ (top). The quadratic coefficient in the latter is statistically significant in *ad hoc* fitting, having magnitude larger than its SE. Note close agreement in slopes (giving E) and in "Chisq" values (sums of squared residuals) for linear fits. C_q values were obtained from log-logistic fits of 24-point regions of profiles centered near the half-intensity points.

$$y = y_0 E^x, \quad (2)$$

where E is the amplification efficiency (AE), ranging from $E = 1$ (no amplification) to $E = 2$ (perfect doubling), x is the cycle number, and y represents the fluorescence signal, or under the assumption that target fluorescence is proportional to its amount, the number N of amplicons. Accordingly, y_0 represents this quantity in cycle 0, before amplification. C_q is defined as the cycle, $x = C_q$, where growth reaches a defined threshold intensity, y_q . Thus a plot of C_q vs $\log(N_0)$ is expected to be linear with slope $-1/\log(E)$. Locating the unknown on this curve by its C_q determines its N_0 .

Although this threshold definition of C_q (often designated C_t) has been most widely used [1,10], it is clear from the plateauing behavior in the growth profiles that Eq. (2) cannot hold throughout the growth region. Recognizing this, most workers have taken y_q as a level near baseline, where it is hoped that Eq. (2) remains valid. However, when the profiles have constant shape, as appears to hold in Fig. 1, the value of y_q is irrelevant, as different choices simply displace C_t by constant amounts [7]. The same holds for other markers, like the first- and second-derivative maxima (FDM and SDM), and Cy0 (the intercept with the cycle axis of a line tangent to the growth curve at the FDM) [10]. The issue then becomes, which of these can be estimated most precisely from the data and also yield a smooth dependence of C_q on $\log(N_0)$. It has been noted that C_b with y_q set near the baseline, actually gives the poorest estimation precision, exacerbated by the practice of separately estimating the baseline level [11,12]. Further, C_t is subject to additional precision loss when the data vary in scale [13,14], as is evident from the varying plateau levels usually observed for the profiles, including those in Fig. 1. The other common markers – FDM, SDM, and Cy0 – are scale-independent, thus free from this problem, as is also the relative threshold C_r , which is obtained by setting y_q to a designated fraction of the plateau level [12].

The least-squares (LS) fit results in Fig. 2 show that indeed the different markers are statistically comparable for the data in Fig. 1. However, the linear response appears to be statistically inferior to quadratic analysis, from which we would conclude that E is concentration dependent. In fact this conclusion is an incorrect consequence of our use of unweighted LS to fit these data, which have excess imprecision at high dilution (low N_0) from unavoidable limitations of Poisson statistics [11]. Such effects can be strong enough to permit quantitative estimation of N_0 directly from the statistical variance in C_q [12,15]. When the data are refitted with appropriate downweighting for the Poisson contributions to the C_q variance (details below), the quadratic contribution is statistically undefined, and $E = 1.916(6)$.

The realization that the slope in Fig. 2 is determined by E has led to efforts to estimate E from analysis of single-reaction (SR) data, in combination with which a single calibration constant relating y_0 to N_0 could permit absolute determination of the unknown N_0 . Most SR methods employ Eq. (2) on data limited to the early growth region [10], while at least two focus directly on y_0 [16,17] but tacitly assume that $E = 2$ in the baseline region [11]. For commonly encountered growth profiles like those in Fig. 1, E must decline from E_0 (~ 2) in the baseline to 1 in the plateau, but there is little direct evidence on just how this decline occurs in the late baseline region.¹ The "mechanistic" models of [16] and [17] predict reasonable decline of E in this region, as does the model [18,19],

$$y(x) = \frac{y_0 y_{\max} E_0^x}{y_0 E_0^x + y_{\max} - y_0} \approx \frac{y_{\max}}{1 + \exp(b(x_{1/2} - x))} \quad (3)$$

where $y_{\max}$ is the limiting growth and E_0 the initial AE. The second expression is the logistic model and is obtained by neglecting y_0 in the denominator of the first; $x_{1/2}$ is the half-intensity point and the FDM, with $b = \ln(E_0)$ and $y_{\max}/y_0 = E_0^{x_{1/2}}$.

When data are generated with the model of Eq. (3) (designated LRE in [19]) and then analyzed using Eq. (2), there is an unavoidable payoff between imprecision and bias: When enough cycles are included to permit adequate precision, E has already declined enough to give significantly low-biased estimates [12]. If there is real variation of E with cycle number, there is another problem with SR methods of estimating it. By simple logic [13], the calibration-based E applies to the

¹ The $E(E_0)$ that appear here are apparent amplification factors, relevant for data analysis. The true AE for the target amplicon may not equal E , especially in the plateau region, where complications from effects like limited dye and amplicon reannealing can lead to a decrease in the fluorescence signal.

earliest cycles: Comparing two starting concentrations, ΔC_q is the number of cycles for N in the more dilute sample to equal N_0 for the more concentrated, after which they grow together. From this, $\log(E) = \log(R)/\Delta C_q$, where R is the dilution ratio. By extension a dilution series estimates E for just the earliest cycles in the most dilute samples to cycle 0 in the most concentrated. Thus, if E varies with cycle number, even a correct estimate of it from early growth cycles could vary substantially from its value in the early baseline cycles.

In the present work we describe an improved algorithm for fitting growth profile data to a log-logistic function with asymmetry parameter [20] plus a baseline function containing 1–4 parameters; and we use it to estimate the 4 common scale-independent C_q markers and their standard errors. We also use its estimates of the SDM and the plateau amplitude to facilitate a second estimation of C_r (which we label $C_{r,x}$), by fitting just cycles up to the SDM to Eq. (2) plus a baseline function [15]. We use these codes to compare the performance of the several markers on six qPCR datasets having replicates at multiple dilutions suitable for calibration fitting, and we ask whether the commonly adopted linear relation is statistically justified. In most cases we find that it is not, implying that E does depend on concentration, and in turn that estimates of N_0 based on linear calibration are biased. The magnitude of such bias may or may not be of practical concern, depending on circumstances. In [1] the authors found that estimates of E can vary with instrument and reaction parameters; here we find that even for a given instrument and fixed reaction conditions, E can vary with N_0 as well as with the choice of C_q marker. These results bear out earlier indications [13] for the large 94×4 Reps technical data set from [10]. Those results were obtained using the C_q estimates provided by the authors of the methods compared in [10]; our present algorithm yields C_q estimates that match or exceed the best of those in their statistical performance.

One observation from our comparisons has important practical significance: Ensemble standard deviations (SDs) for C_q typically exceed the LS parametric standard errors (SEs) by factors of 2–5. The SEs are based on the fit model and the random error in the profile data, and they should correctly predict the ensemble SDs if such data error is the only source of dispersion in the estimates [12,21]. The excess dispersion in the ensemble estimates thus means there are other sources of experimental error, which likely include pipette volume uncertainty. It follows that qPCR is capable of much higher precision through better experimental procedures.

2. Mathematical background and methods

2.1. Amplification efficiency from calibration fitting

Fitting data to a straight-line response function is perhaps the most common exercise in data analysis and needs little review. Here the fit relation is

$$C_q = a + bx = a - x/\log(E), \quad (4)$$

where $x = \log(N_0)$. The second expression in Eq. (4) represents a non-linear LS (NLS) fit, with the advantage of yielding the parametric standard error (SE) directly from NLS programs that provide SEs. Alternatively, the SE in E can be obtained from that in b using error propagation [11,22],

$$\sigma_E = E \ln(10) \sigma_b/b^2. \quad (5)$$

The extension of Eq. (4) to polynomials of order 2 and higher is straightforward and remains a linear LS fit,

$$C_q = a + bx + cx^2 + \dots, \quad (6)$$

However, the slope is now a function of x , involving b and all higher-order fit parameters. On the other hand, by recentering the fit about x_0 (= selected $\log(N_0)$), we obtain a statistically equivalent fit

where b is the slope at x_0 [22], so E and its SE can be obtained as a function of x_0 by just changing x_0 in the expression,

$$C_q = a + b(x - x_0) + c(x - x_0)^2 + \dots = a - (x - x_0)/\log(E) + c(x - x_0)^2 + \dots, \quad (7)$$

Calibration fitting in qPCR is almost universally done with neglect of weighting, which tacitly assumes the data have constant uncertainty. This assumption is incorrect when the calibration data extend to N_0 small enough to make Poisson uncertainty in N_0 significant – typically $N_0 \approx 100$ or smaller [13,15] – in which case weighted fitting is called for. As is described below, one can often get reliable estimates of the C_q variance from all other effects by analyzing the replicates at the higher concentrations. The Poisson contribution for small N_0 is approximately

$$\sigma_{C_q, \text{Pois}}^2 = [(\ln(E))^2 N_0]^{-1}, \quad (8)$$

and can be added to the estimated variance from other effects to give the total variance. The weights w_i are then taken as the reciprocals of the total variances.

2.2. The log-logistic model with asymmetry

The log-logistic model is a sigmoidal alternative to Eq. (3), and in the 4-parameter form [20],

$$\text{LL4}(x) = y_{\max} [1 + (g/x)^h]^{-p}, \quad (9)$$

it can accommodate some asymmetry. With $p = 1$, it is nearly symmetrical, and $g \approx x_{1/2}$ ($\approx x_{\text{FDM}}$). We have found that this form, in combination with suitable baseline functions, can yield C_q values that are statistically at least as precise as those values obtained by any of the other methods reviewed in [10]. We achieve this performance by fitting typically 20–30 cycles in the transition region; and we evaluate all 4 common C_q markers at a time: FDM, SDM, Cy0, and relative threshold C_r , all of which are scale-independent.

The FDM is obtained by setting the second derivative of $\text{LL4}(x) = 0$ and solving for $x = x_{\text{FDM}}$. Similarly, the SDM is obtained by setting the third derivative = 0. Cy0 is defined [23] as the intercept with the cycle axis of a straight line tangent to the growth curve at the FDM. C_r is obtained by first estimating the plateau level $y_{\max}$ and then solving $\text{LL4}(C_r) = r y_{\max}$, with r a specified fraction, typically about 0.15. In fitting data to Eq. (9), there is always at least one additional parameter for the baseline, and we have used as many as four, giving from 5 to 8 adjustable parameters (see below).

The expression for the FDM is

$$x_{\text{FDM}} = g/Z^{(1/h)}, \quad (10)$$

where $Z = (h + 1)/(hp - 1)$. We can obtain the FDM and its SE directly from the fit by replacing g with x_{FDM} , whereupon the denominator in Eq. (9) becomes D^p with

$$D = 1 + \left(\frac{h + 1}{hp - 1}\right) \left(\frac{x_{\text{FDM}}}{x}\right)^h. \quad (11)$$

This substitution can be useful, because x_{FDM} is both more precise and easier to estimate visually than g for curves with significant asymmetry. Cy0 is given by

$$\text{Cy0} = x_{\text{FDM}} \left(1 - \frac{(1 + Z)}{hpZ}\right), \quad (12)$$

and the SDM is obtained by solving the quadratic equation in u ($= (g/x_{\text{SDM}})^h$),

$$Au^2 + Bu + C = 0, \quad (13)$$

with $A = (hp - 1)(hp - 2)$, $B = (h + 1)(4 - h - 3hp)$, and $C = (h + 1)(h + 2)$.

For asymmetrical growth profiles, there are actually two different modes on which the fits can converge: with $h > 0$ and $h < 0$. Correspondingly, the roles of the parameters for baseline and plateau

are reversed, with $y_{\max}$ becoming negative; and p changes from > 1 to $0 < p < 1$. The two modes converge for $p = 1$ but are not equivalent otherwise, typically showing a factor of ~ 2 difference in fit variance. We examine both modes below. Note that there are no parameters in this model for predicting E_0 in the baseline region, and in fact the calculated values of $E(x)$, from $LL4(x+1)/LL4(x)$, are not physically reasonable outside of the growth region. Thus, estimating E_0 with the LL4 model in SR analysis requires some prescription for choosing the appropriate x at which to assess this ratio. We do not pursue this matter further in this work.

As mentioned earlier, we have found that the C_q markers are obtained with optimal precision by limiting the fit to just 20–30 cycles centered on the growth region. We designate the fit region by first locating the approximate SDM (SDM_{prx}) from second differences, as done by Boggy and Woolf [16], and then specifying a start cycle $x_1 = SDM_{\text{prx}} - \Delta_1$ and end cycle $x_2 = SDM_{\text{prx}} + \Delta_2$. By trial and error, we find best results with $\Delta_1 = 8$ –12 and $\Delta_2 = 12$ –16, depending on dataset.² The C_q estimates do depend somewhat on x_1 and x_2 , and in cases where the true SDM falls near a half-cycle, SDM_{prx} can vary with reaction in a set of replicates, giving excess C_q variance from the variation in cycle range. To reduce such effects, we can use code versions that permit specifying absolute x_1 and x_2 , which would normally be employed after discovering results that show varying ranges among replicates. With all replicates analyzed with the same x_1 and x_2 , the precision is not strongly sensitive to these values, with variance typically changing by only a few percent when x_1 and x_2 are changed by ± 1 .

By using the LL4 estimates of $y_{\max}$ and SDM, we are able to estimate C_r a second way: fit just cycles up to the (rounded) SDM to the exponential growth law of Eq. (2) plus a suitable baseline function (discussed below). This threshold, which we label $C_{r,x}$, is a 1-step relative version of the FPLM method [10,24] and was used previously to estimate absolute copy number from the C_q variance [15].

2.3. Baseline functions

It is common practice to estimate the baseline from a selected range of early cycles and then subtract it from the profile to yield the growth curve. As compared with simultaneous estimation of baseline and growth parameters from a single nonlinear LS fit, this 2-step procedure is statistically inferior, yielding biased C_q estimates that on average amount to a 3-fold increase in C_q variance [12]. The 1-step NLS fit is easy to implement, with

$$y(x) = \text{bas}(x) + LL4(x), \quad (14)$$

being the fit model for the LL4 growth curve. The choice of function for $\text{bas}(x)$ can depend on the range of early cycles included in the fit. We have commonly used linear and quadratic functions of x , in which we judge the need for parameters beyond the minimal single constant by their statistical significance in the fit: Any parameter having SE greater than its magnitude is statistically undefined in *ad hoc* fitting [25]. For data like those in the 94 $\times$ 4 Reps dataset from [10] that exhibit "saturation" baseline behavior, we have used the exponential plateau function [10,15],

$$\text{bas}(x) = a - q \exp(-px), \quad (15)$$

sometimes with an added linear term. Saturation behavior manifests as a rise in the first few cycles, so if these are omitted from the fit range, a linear or quadratic $\text{bas}(x)$ may perform as well.

² This specification can as well be tied to the approximate FDM from first differences, in which case Δ_1 and Δ_2 increase and decrease by ~ 2 cycles, respectively, giving a range that is approximately symmetric about the FDM.

2.4. Weighted least squares

In linear LS it is rigorously true (but not widely appreciated) that minimum-variance parameter estimates are obtained if and only if the data are weighted inversely as their variances. The same is normally assumed to hold in nonlinear LS but cannot be proved, in part because many NLS estimators do not even have finite variance. Still, Monte Carlo simulations for common nonlinear models have supported inverse-variance weighting [21,26]. This issue arises in the LS fitting of qPCR data in two situations already mentioned: the analysis of reaction profile data by whole-curve fitting, and the fitting of C_q vs $\log(N_0)$ for calibration [13].

Considering first the fitting of growth profile data to sigmoidal models, most such data are collected on instruments that monitor fluorescence. A major source of random error or noise in optical data has long been referred to as "flicker" [27], describing the fluctuations in the light source. The result is noise proportional to signal; for example, 1% fluctuation in the intensity of the light source produces 1% fluctuation in the detected signal ($\sigma_{yp} = \sigma_p y$). However, as the signal decreases, the limiting noise becomes a constant related to the properties of the detection instrumentation. The sources are considered independent, so their variances add, giving a simple relation that holds well for a number of different experimental techniques [28,29]:

$$\sigma_y^2 = \sigma_0^2 + (\sigma_p y)^2. \quad (16)$$

The main effects of neglecting weights are reduced precision of estimation and incorrect parametric SEs; however, the magnitude of such losses may be tolerable when the weights vary by less than a factor of 10 over the data set [8,30]. This, for example, appears to be the case for the 94 $\times$ 4 Reps data from [10], where the baseline intensity level is about half that in the plateau region. However, it does not hold for most qPCR data we have analyzed, and is borderline for those in Fig. 1, where plateau levels are about 3 times baseline levels.

The conditions requiring weighting in calibration fitting of C_q vs $\log(N_0)$ were discussed in Section 2.1. When the range of N_0 includes values small enough to add Poisson scatter to the results, the increased variance dictates a reduced weight for the relevant N_0 values. The Poisson variance contribution is predictable and was given in Eq. (8).

In assessing the quality of LS fits, an important metric is χ^2 [13], defined as the sum of weighted squared residuals, $\sum w_i \delta_i^2$, where δ_i is the difference between observed and calculated y_i . If the data variances σ_{yi}^2 are known and the w_i are taken as their inverses, the expected value of χ^2 is the number of statistical degrees of freedom $\nu = n - p$, where n is the number of fitted values and p the number of adjustable parameters. Accordingly the expected value for the reduced χ^2 (RCS, χ^2/ν) is 1. The standard deviation (SD) of the RCS is $(2/\nu)^{1/2}$, so that, e.g., the 90% range is 0.39–1.83 for $\nu = 10$, narrowing to 0.66–1.39 for $\nu = 40$ [25]. Too-high values of RCS result from some combination of optimistic assessment of the data error and a poor fit model, while too-low mean pessimistic data errors. In unweighted fitting the data are tacitly assumed to have constant variance and the w_i are taken as 1; RCS then becomes an estimate of the data variance σ_y^2 . When comparing fit models, smaller χ^2 indicates a better fit.

2.5. Computations

We have used the KaleidaGraph program (Synergy Software) for examining data and running preliminary analyses, also for preparing the illustrations in this work. For processing the multireplicate data sets efficiently, we have devised FORTRAN (Microsoft) codes that are similar in structure to the NLS routine provided long ago by Bevington (program 11-5, CURFIT) [25]. The different C_q markers are evaluated as detailed above, and their SEs are obtained by error propagation [22]. The NLS fits require initial values for the parameters but are not very sensitive to these, making it possible to achieve successful convergence

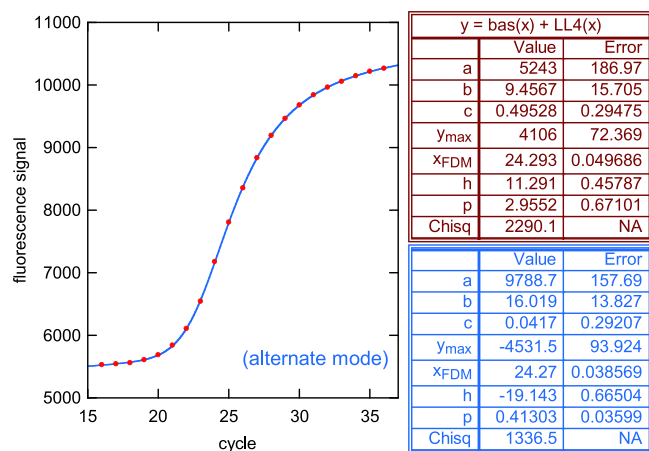


Fig. 3. NLS fits of first reaction at highest concentration in 94 × 4 Reps dataset [10] to LL4 + bas(x) = $a + bx + cx^2$, in normal (upper) and alternate (alt) modes. The quantity D in the denominator of LL4 is as defined in Eq (11). Chisq is the sum of squared residuals for these unweighted fits.

for all but at most a few experiments in a multireplicate data set in a first attempt, with success on problem experiments in a second pass.

User-friendly versions of these programs are planned for later distribution; similar codes in R are already available [20], and modifications that permit most of the computations described here can be obtained at <https://github.com/anspiess/qPCR-algorithms>.

3. Results and discussion

3.1. Performance tests

Fig. 3 illustrates the two convergence modes for the LL4 model on one of the reactions in the 94 × 4 Reps dataset from [10]. Note the negative values for $y_{\max}$ and h in alt mode and the reversed role for a from baseline to plateau level. For this example, the slope b is statistically insignificant in normal mode and the quadratic coefficient is insignificant in alt mode. The significantly smaller Chisq (χ^2) value in alt mode is characteristic of all reactions in this dataset.³ As has been noted above, unweighted NLS suffices for these fits, because the data noise varies by only a factor of ~ 2 from baseline to plateau. (Other reasons are discussed below.) The close agreement in the two estimates of the FDM holds in general for all C_q markers except the SDM, where there is systematic difference of about 0.2 (alt mode higher). This effect and the different χ^2 values presumably relate to differences in how the asymmetry is handled in the early growth region vs the approach to plateau in the two modes.

In [13] (Table 1 and Fig. 3) were compared the precisions of estimation of C_q for each of the 4 concentrations in this dataset for the 7 methods reviewed in [10]. The minimum-variance results came from the Miner method [31] for 3 of the 4 concentrations, with Cy0 slightly bettering it for the most dilute samples. The closest to Miner was 5PSM [20], in which the LL4 model was employed but with a single baseline parameter and fitting all cycles. As we have noted, the present version of this model gives ensemble precisions equaling or bettering the best in [10]. Cy0 and C_r gave generally smaller ensemble SDs across the concentration range than FDM and SDM, with the two modes being comparable in all cases. Fig. 4 compares the Cy0 ensemble SDs with the smallest from [10] and with $C_{r,x}$ estimates obtained in [15] using Eq. (15) for the saturation baseline. Also shown in Fig. 4 are the rms averages of the SEs predicted by the individual fits for Cy0 in alt mode

and for the $C_{r,x}$ model. The point of this comparison is to emphasize that the ensemble SDs display excess dispersion from effects other than data noise. If the latter were the only source of variability, we would expect the ensemble SDs to agree with the parametric SEs [21]. As already noted, the excess dispersion at small N_0 is from Poisson scatter. At large N_0 Poisson scatter is negligible, and we have suggested that pipette delivery volume error is the main source of excess dispersion [12,15]. Reasons for the much smaller predicted SE for the $C_{r,x}$ model are discussed just below. The better fit quality for LL4 in alt mode is responsible for its predicted SEs being $\sim 25\%$ smaller than those for normal mode.

The observation that the ensemble SDs significantly exceed the parametric SEs means that the LL4 model used to estimate the C_q markers is likely adequate for this task, thanks to the excess dispersion from sources other than data noise. However, neither the LL4 model nor a similar 4-parameter version of Eq. (3) accurately represents the fitted data, as is clear from the pronounced systematic component in the fit residuals (see graphic and Supplemental Fig. S-5 in [12]). The improved fit quality in alt mode reduces but does not eliminate these effects. On the other hand, systematic residual trends practically vanish for the $C_{r,x}$ model, because only a few cycles in the growth region are included in the fits. This leads to an estimated σ_y a factor of 3 smaller than that from the LL4 fits, and this is a major contributor to the smaller parametric SE for this model in Fig. 4. Poisson dispersion for small N_0 is an unavoidable mathematical reality; however, the excess dispersion at large N_0 in these data is presumably reducible through better experimental techniques, in which case better fit models will be needed to take advantage of the improved precision in whole-curve fitting. The low rms SE for $C_{r,x}$ represents the best that can be hoped for optimal experimental data; its factor-of-4 improvement over the ensemble SDs at large N_0 represents a 4^2 efficiency improvement, meaning one reaction would be the statistical equivalent of 16 replicates.

As we have noted, one source of excess dispersion is correctable through data analysis alone, namely the effect of scale variability on the absolute threshold C_t . Fig. 5 compares the ensemble SDs for C_t and $C_{r,x}$ with y_q and r chosen to make the average C_q s about the same. For the largest two N_0 s, the variance ratios are 2 and 5; these would have been even larger had we used the standard practice of separately estimating the baseline and subtracting before assessing C_t [12]. The large differences stem directly from the pronounced scale variability for these data [13,14].

Fig. 6 shows calibration fit results for this dataset. The weights for all markers at each concentration were the same as used to produce Fig. 5 in [13], so the low Chisq values (cf $\nu \approx 370$) are another indicator of the higher precision of the present C_q estimates. The results resemble those in Fig. 4 from [13] in showing a statistically insignificant cubic coefficient for C_r . This parameter was also insignificant for FDM, was barely so for Cy0, and significant by about 2σ for SDM. Results for C_q from alt mode were very similar, with Chisq smaller by as much as ~ 25 for SDM, higher by ~ 5 for FDM. Fig. 7 shows that the E estimates from the quadratic fits are statistically consistent for all markers. The increase in E with concentration is solid, but at the highest concentration, all E s exceed the physical limit of 2.0 by more than 1σ .

3.2. Weighted fitting

The need for weighted LS arises in two situations, already noted: In the estimation of C_q by fitting whole-curve data when the scatter in the plateau region greatly exceeds that in the baseline region; and in calibration fitting when the C_q range includes N_0 values small enough (< 100) that Poisson scatter contributes significantly to the C_q variance. Using the data behind Figs. 1 and 2, we show how to derive the needed weights.

Fig. 8 illustrates the estimation of the variance in the baseline and plateau regions for the curves shown in Fig. 1. The approach is to use the scatter about a smooth approximation of the data in regions where

³ This behavior is not general. An examination of representative data from the profiles shown in Fig. 1 and from two additional datasets discussed further below showed χ^2 for alt mode higher in all 16 cases, by factors from 1.2–4.4.

Table 1

C_q values obtained for 3×5 data from Rutledge and Stewart [7] by fitting to the LL4 model with a linear baseline function (Eq. (14)). Cycles in the indicated range (c_1 – c_2) were fitted, with values inverse-variance weighted using the variance function from Fig. 9.^a

N_0	c_1 – c_2	RCS ^b	FDM	SDM	C_r ^c	Cy0	$C_{r,x}$ ^{c,d}
188,000	7–30	0.826	18.408 (22)	16.420 (28)	15.549 (23)	15.457 (20)	15.442 (30)
	7–30	1.325	18.447 (28)	16.493 (37)	15.525 (30)	15.482 (25)	15.373 (38)
	7–30	0.96	18.444 (24)	16.497 (31)	15.518 (25)	15.481 (21)	15.402 (38)
18,800	10–33	1.359	22.014 (28)	20.052 (37)	19.066 (31)	19.030 (25)	18.972 (41)
	11–34	1.013	21.983 (24)	20.025 (31)	19.056 (27)	19.014 (23)	18.984 (27)
	11–34	1.16	21.972 (26)	20.011 (34)	19.112 (29)	19.039 (25)	18.934 (27)
1880	15–38	0.845	25.408 (22)	23.441 (29)	22.515 (25)	22.454 (22)	22.437 (29)
	14–37	1.028	25.553 (24)	23.587 (31)	22.609 (27)	22.570 (22)	22.544 (27)
	15–38	0.752	25.470 (20)	23.528 (27)	22.562 (23)	22.523 (20)	22.422 (27)
188	18–41	1.741	29.196 (31)	27.232 (40)	26.306 (35)	26.247 (28)	26.211 (21)
	18–41	1.501	29.073 (29)	27.109 (38)	26.160 (34)	26.111 (28)	26.097 (23)
	19–42	1.41	29.376 (28)	27.376 (37)	26.506 (32)	26.415 (28)	26.471 (37)
18.8	21–44	1.252	32.407 (27)	30.455 (35)	29.558 (31)	29.491 (25)	29.514 (28)
	21–44	1.558	31.987 (31)	30.040 (40)	29.165 (36)	29.090 (29)	29.028 (26)
	22–45	1.14	32.530 (25)	30.585 (34)	29.687 (30)	29.622 (25)	29.558 (19)

^a Figures in parentheses are parametric standard errors, in terms of final displayed digits; e.g., 18.408 (22) means SE = 0.022.

^b Reduced chi-square.

^c Obtained for $y_q = 0.12 y_{\max}$.

^d Obtained fitting cycles 5– c_2 to Eq. (2) plus a linear baseline, with $c_2 = 15, 19, 23, 26, 29$ for the 5 concentrations, respectively.

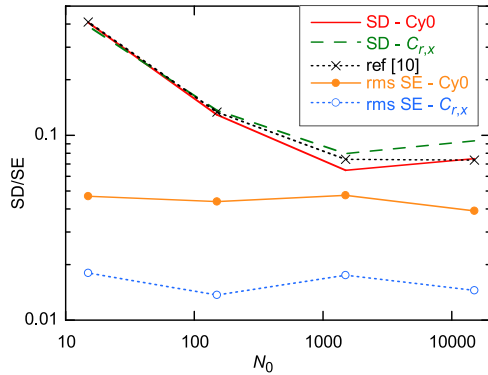


Fig. 4. Standard deviations/errors for each of 4 concentrations in the 94×4 Reps data [10]. Ensemble SDs at top from present estimates of Cy0, compared with best from [10] and $C_{r,x}$ estimates from [15]. At bottom are the rms (root-mean-square) averages of the parametric SEs from the individual fits, for Cy0 using LL4 model in alt mode, and for $C_{r,x}$. Connecting lines are just for display purposes.

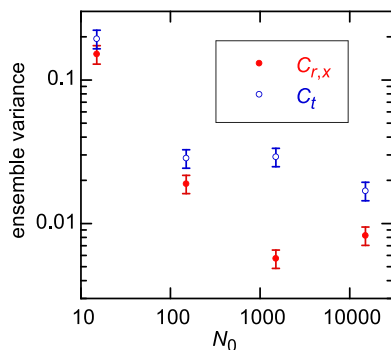


Fig. 5. Ensemble variances for absolute and relative threshold in the 94×4 Reps data. For $C_b, y_q = 700$; for $C_{r,x}, r = 0.18$. Estimates for both were obtained by fitting to Eq. (2) plus the $\text{bas}(x)$ function of Eq. (15). Error bars represent one SD. The average C_t values slightly exceed those for $C_{r,x}$, by from 0.07 to 0.30.

they vary little in magnitude. For that approximation, here we use polynomials of 4th order fitted by unweighted LS. From the discussion near the end of Section 2.4, the estimated variance is Chisq/ν , with

$\nu = n - 5$, giving estimated $\sigma_y^2 = 173$ and 1845, respectively, for the baseline and plateau regions from the included fit results. Note that all polynomial parameters are statistically significant here for the baseline fit (B), since all SEs are less than the parameter magnitudes; however, this is not true for the plateau data, and the indicated reduction of the fit order decreases that variance estimate to 1554.

Fig. 9 shows the fit of the 15 baseline and 15 plateau variance estimates from these data to Eq. (16). This is a weighted fit, since variance estimates have error proportional to $(2/\nu)^{1/2}$, i.e., $\sigma(\sigma_y^2) = (2/\nu)^{1/2} \sigma_y^2$. The resulting χ^2 is statistically reasonable –38 as compared with $\nu = 28$. The data here are not sufficient to confirm that Eq. (16) is better than other variance functions (VF); however, this function works well in many situations and VFs are anyway not needed with great precision to yield near-optimal fit results when they are subsequently used in weighting. An alternative to the fit shown in Fig. 9 is the fit of $\ln(\sigma_y^2)$ to $\ln(a^2 + (by)^2)$, in which the weighting is simpler [28]: $\sigma(\ln(\sigma_y^2)) = (2/\nu)^{1/2}$. Since ν varies with estimate here, weighted fitting is still required.

Using this VF to compute weights (inverse variances), we fitted the 15 profiles to Eq. (14), obtaining results for the 5 C_q markers summarized in Table 1. To do the calibration analysis of Fig. 2 properly, we need to include weights for C_q to compensate for the decreasing precision of these with decreasing N_0 from Poisson scatter. Variances estimated from just 3 values are inherently very uncertain: $(2/\nu)^{1/2} = 1$, meaning 100% relative SD. However, using the estimates collectively, we obtain adequate results by assuming the C_q variance is the sum of the predictable Poisson contribution and a constant, as shown in Fig. 10. When weights computed from these variance functions are used in the calibration fits of Fig. 2, there is no statistically significant nonlinearity for any of these C_q markers; and they yield virtually identical E : 1.918 for C_r , 1.916 for the others.

3.3. Other multireplicate datasets

In addition to the two datasets already discussed, from [7] and [10], we have analyzed data from Rutledge and Cote [5] (20 replicates at each of 6 concentrations), Guescini et al. [23] (12 reps at 7 concentrations), Lievens, et al. [32] (18×5), and Karlen et al. [33] (FN, $15 \times 4 + 12$). Of these, the Rutledge and Cote data showed only marginally significant quadratic coefficients for just two of the 5 C_q markers, leading to an increase in E from 1.90(1) to 1.93(1) from the most to least concentrated samples. Linear fits gave E s spanning the narrow

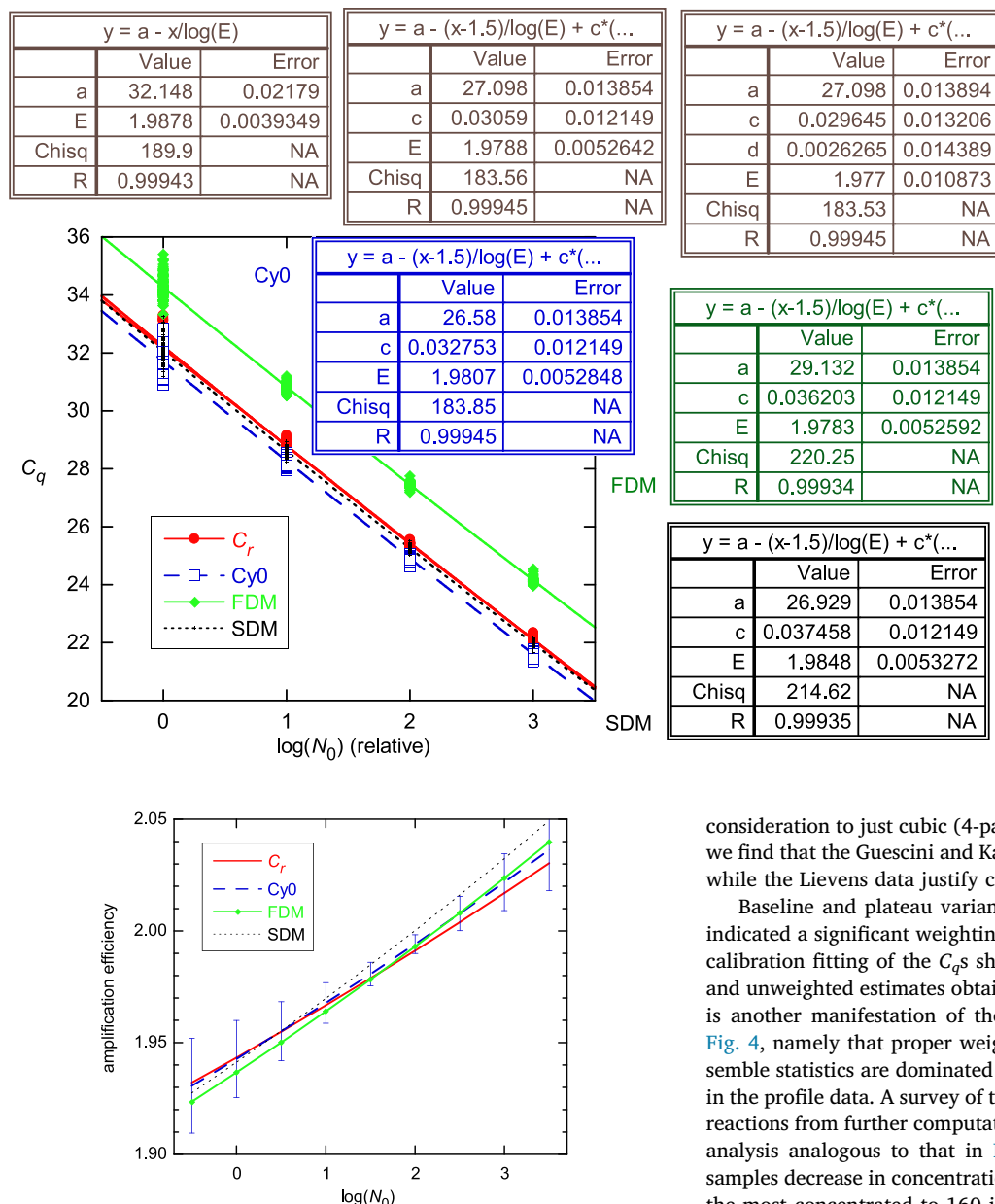


Fig. 7. Amplification efficiency as a function of concentration, from quadratic fit results in Fig. 6. Error bars (1- σ) are shown for just Cy0 but are nearly identical for all 4 markers.

range 1.902(3) to 1.916(3). As discussed below (and see online Supplement for more detail), the other three datasets showed statistically significant nonlinearity in calibration fits. The data themselves, some of which are no longer downloadable from the journal web sites, can be obtained at www.dr-spiess.de.

For the other three datasets, at least some C_q markers gave statistical significance for all 5 parameters in a fit to a quartic polynomial, and the Lievens data did so for all 5 markers. However, drawing reliable conclusions about the derivatives (hence E) from such high-order fits is risky, as the fit functions typically diverge from the fitted points rapidly outside their range. Also, the polynomial representation is just a means to an end, and typically other 4- and 5-parameter functions can give comparably small χ^2 and yet yield substantially different AEs. Quantifying such "model error" requires examining other functions in trial-and-error fashion, a possibly demanding task. Here our purpose is less the precise determination of E than seeing whether it varies with N_0 and how that might affect calibration results. Restricting our

Fig. 6. Calibration fits of C_q estimates for 94×4 Repts data, weighted using a common set of inverse ensemble variances. At top are linear, quadratic, and cubic fits of the C_r estimates to polynomials in $(x-1.5)$, showing that the cubic coefficient (d) is not statistically defined but that the quadratic one (c) is. For comparison, the quadratic fits of $Cy0$, FDM, and SDM are included, confirming that c is statistically significant in every case and showing that all E estimates are statistically consistent at $x = 1.5$.

consideration to just cubic (4-parameter) and lower-order polynomials, we find that the Guescini and Karlen data require quadratic calibration, while the Lievens data justify cubic representation.

Baseline and plateau variance estimates for the Lievens data [32] indicated a significant weighting range (~ 500), but in fact subsequent calibration fitting of the C_q s showed little difference for the weighted and unweighted estimates obtained by fitting with Eq. (14). This result is another manifestation of the effects discussed in connection with Fig. 4, namely that proper weighting doesn't help much when the ensemble statistics are dominated by factors other than the random noise in the profile data. A survey of the C_q estimates led to the exclusion of 8 reactions from further computations (see Supplement). The C_q variance analysis analogous to that in Fig. 10 is shown in Fig. 11. Here the samples decrease in concentration by factors of 5 from 10^5 templates in the most concentrated to 160 in the least. However, at low N_0 the C_q variances for all markers were higher than predicted, so we added a correction factor for N_0 in the fit model, yielding values ranging from 0.37 for SDM to 0.61 for $C_{r,x}$. These values imply that N_0 in the most dilute sample is 59–98, in rough agreement with the results from a similar analysis in [15]. To obtain a common variance function for all C_q markers, we set the correction factor at its average (0.42) and refitted, obtaining an average for the constant $A = 0.0005$. This VF was then used to weight the C_q s in the calibration fitting, yielding the cubic-based E results illustrated in Fig. 12. Although the case for N_0 -dependence in E is solid, the several C_q markers give AEs that often differ by more than their combined SEs. This statistical inconsistency can be taken as a rough indicator of model uncertainty, since the fit quality (χ^2) does not vary strongly.

The ranges of data error in the growth curves for the Guescini [23] and Karlen FN data [33] were about 30 and 100, respectively, or enough to warrant weighting in the fits to extract C_q . However, here again visual surveys of the C_q estimates showed patterned variation that greatly exceeded the differences between weighted and unweighted estimates, indicating that the ensemble C_q variances are dominated by effects other than data noise (see Supplement). Indeed, for the Karlen data, the highest C_q variances occurred for the most concentrated

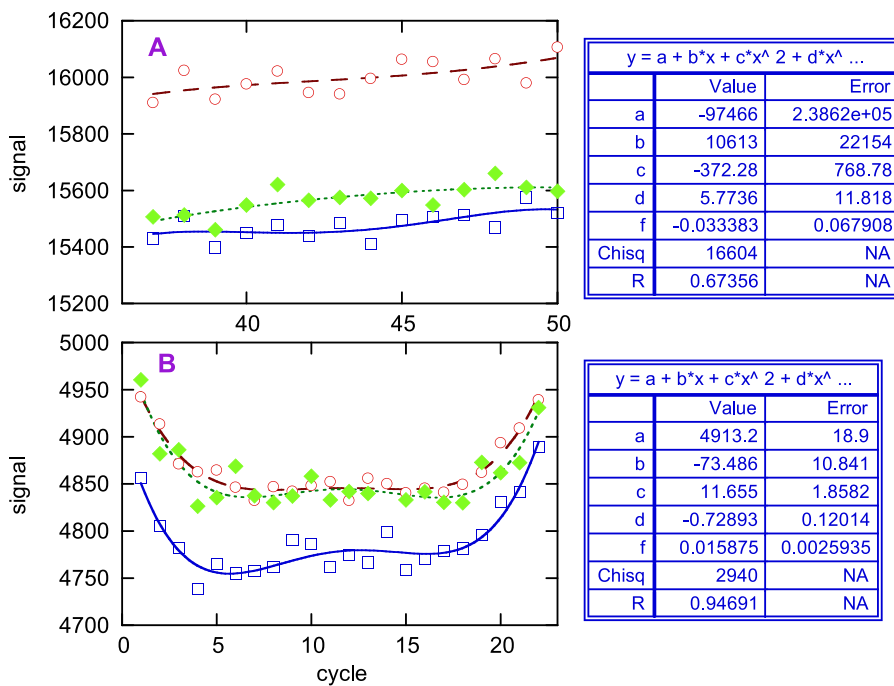


Fig. 8. Estimating data variance from polynomial fitting, for 4th dilution in 3×5 data from [7], in plateau (A) and baseline (B) regions. The estimated variances are $\text{Chisq}/(n-5)$, with $n = 14$ in the plateau region and 22 in the baseline region. Fit results are shown for only the lowest curve in each panel; Chisq values for the other curves (open and solid points, respectively) are 27,000 and 14,700 (A) and 1056 and 4080 (B). Note that none of the parameters in A is statistically significant; in fact these data are well represented by a quadratic function, with little increase in Chisq but an increase of 2 in ν , giving $\sim 20\%$ smaller estimated variances.

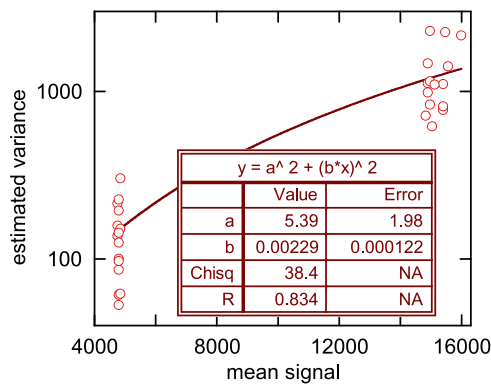


Fig. 9. Fit of estimated variances for 3×5 data from [7] to Eq. (16). From these results, the second term dominates the variance even in the baseline region.

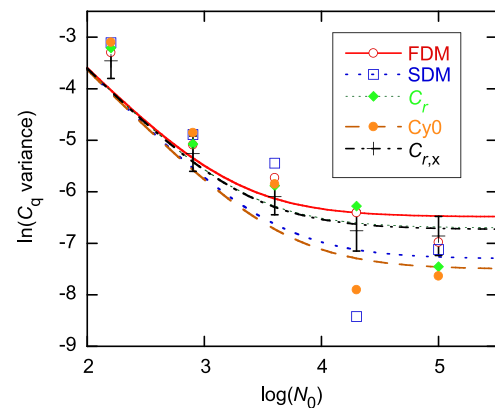


Fig. 11. C_q variance estimates from Lievens data [32], displayed and fitted as in Fig. 10, with E taken as 1.86 and $N_0 = 160$ for the lowest concentration. Error bars shown for $C_{r,x}$ are representative of the others.

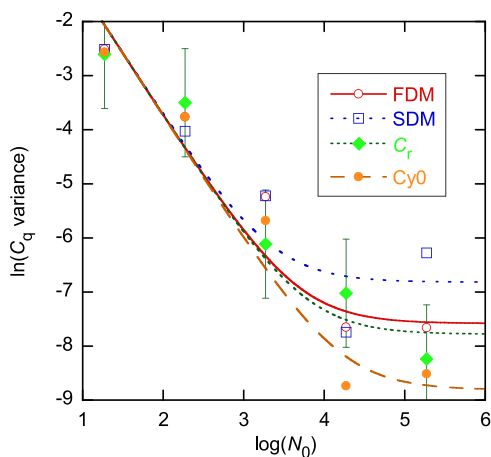


Fig. 10. C_q variance estimates from replicate values in Table 1, displayed in logarithmic form, and results from fitting values for each marker to $\ln(A + \sigma^2 C_{q,\text{Pois}})$, where the Poisson variance is given by Eq. (8), with E taken as 1.915. Error bars are shown for C_r only but are the same for all, $\sigma = 1$. Values of A range from 0.00015 for Cy0 to 0.0011 for SDM.

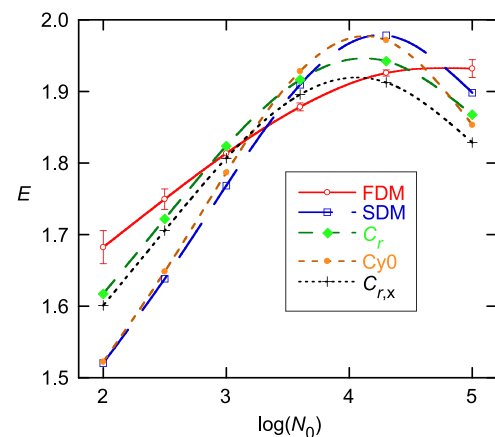


Fig. 12. Dependence of AE on N_0 from cubic calibration fits of C_q estimates for data from [32]. Error bars shown for FDM are comparable for all. χ^2 in the calibration fits ranged from 87 for Cy0 to 112 for SDM (90 C_q values).

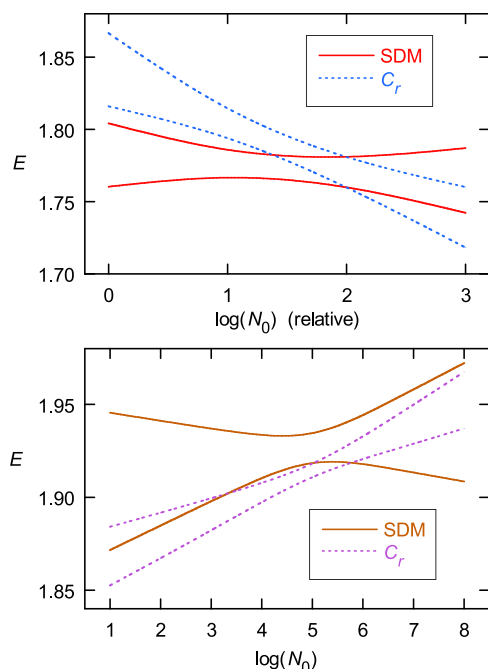


Fig. 13. 1- σ confidence bands for extreme estimates of E as functions of $\log(N_0)$ for data from [23] (lower) and [33] (upper). For the former (84 reactions), χ^2 values for FDM, SDM, C_r , Cy0, and $C_{r,x}$ were, respectively, 100, 500, 109, 183, and 109; in the same order the χ^2 values for the latter (72 reactions) were 80, 79, 62, 66, and 71.

samples, where Poisson scatter is negligible. Thus these data were weighted in the calibration fits using just ensemble variances in place of the approach of Figs. 10 and 11 (which was used for the Guescini C_q s). In both cases a single average C_q variance function was used for all markers, to permit comparison of their fit qualities through their χ^2 values.

As already indicated, both datasets support quadratic calibration fits, giving slopes linear in $\log(N_0)$ (but nonlinear AEs, thanks to their inverse slope exponential dependence). The extreme E estimates are illustrated in the form of error bands in Fig. 13. In both cases the SDM estimates (and *only* these) actually support constant E , though for the Guescini data, this is because of the anomalously poor quality of the calibration fit ($\chi^2 = 500$). In contrast, the C_r fits give lowest or near-lowest χ^2 values and show clear dependence of E on N_0 .

For these 4 datasets, the precision comparisons resemble those in Fig. 4, with the greatest disparity between rms parametric SEs and ensemble SDs for the Karlen data (ratio ~ 10) and the least for Rutledge and Cote (~ 2). The SEs are compared in Fig. 14, which shows that $C_{r,x}$ is lowest except for the Lievens data, where Cy0 is slightly lower. We are reluctant to interpret these data any further, as some have clearly been baselined and may also have been smoothed. For statistical comparisons like those we are attempting here to be valid, fully "raw" data are needed [12,34].

3.4. E is not constant; so what?

The consequences of nonconstant E depend on whether and how this E is used. Thus, if calibration data are fitted to a linear response (assuming constant E) instead of a quadratic, E is not needed explicitly and the resulting bias is likely to be minor as long as the unknown is in the range of the calibration data. For example, with C_r calibration of the Guescini data (Fig. 13), the greatest in-range bias from linear calibration (which gives $E = 1.913$) is only -0.04 in $\log(N_0)$ for the unknown, which is about half the magnitude of the SE for a single measurement, and which translates into a 9% undershoot in the estimate of N_0 . Since

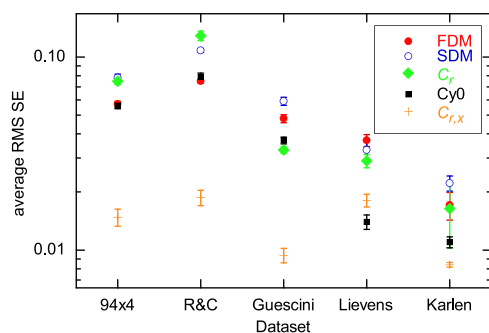


Fig. 14. rms parametric SE values for the different C_q markers, averaged over all concentrations in each dataset. These SEs generally vary little with concentration.

replicates reduce the SE roughly as $n^{-1/2}$, this bias does become marginally significant for 4 replicates.⁴ Alternatively, if E is used to estimate N_0 for an unknown relative to a single known (as is the hope behind SR methods), the error increases with the difference ΔC_q for the reference and unknown. For example a 5% error in E gives a factor $1.05^{\Delta C_q}$ (or its reciprocal), which is 2 (or 1/2) for $|\Delta C_q| = 14$. (The error is smaller if 5% is the maximum error in variable E , as for the Guescini data in Fig. 13). (See Supplement for a fuller discussion of this matter.)

The significant dependence of estimated E on choice of C_q marker for the last three datasets considered here means that their growth curves are changing shape systematically with N_0 . When such changes occur smoothly, there is little effect on calibration results, an exception being the SDM for the Guescini data, which gave anomalously large χ^2 . For best estimates of the "true" E , the threshold $C_{r,x}$ might be least sensitive to such N_0 -dependent shape changes. Note that even though this method yields low-biased SR estimates of E [12], its C_q estimates remain valid for calibration fitting. Note further that the whole-curve-based threshold marker C_r gave comparable χ^2 in calibration fitting and E s that were statistically consistent with those from $C_{r,x}$.

4. Conclusion

A nonlinear LS algorithm for fitting qPCR growth profiles to a 4-parameter log-logistic function [20] has been modified to incorporate multiparameter baseline functions and permit fitting selected cycle ranges, with output of 5 common scale-independent C_q markers and their SEs. Statistical comparisons for the very large 94×4 Reps dataset show that the C_q estimates from this routine are comparable to the best obtainable by other methods. The provision for limiting the fitting to selected cycle ranges means that whole-curve models can be used even on data with anomalous profiles, like the "hook" behavior sometimes seen in the plateau region [35].

The dependence of C_q on $\log(N_0)$ shows statistically significant nonlinear behavior for 4 of 6 multireplicate datasets, which means that the AE is not constant for these. When such data are used in calibration, neglect of this nonlinearity leads to bias in the estimates of N_0 ; however, this bias is likely to be tolerably small in most applications. On the other hand when linear-based E estimates are used in lengthy extrapolations from a reference sample, the bias can become significant. This could be important when E is estimated from single-reaction data, because such estimates (1) focus on the early growth cycles, and (2) typically are low-biased. Even without the bias, such estimates are perforce constant and may not well approximate E in the earliest baseline

⁴ The SE for the unknown includes contributions from the calibration curve, but here that is small compared with a single measurement of the unknown, thanks to the large number of replicates. In general, optimal calibration is achieved with about equal numbers of knowns and unknowns, requiring both to be increased together to best increase precision.

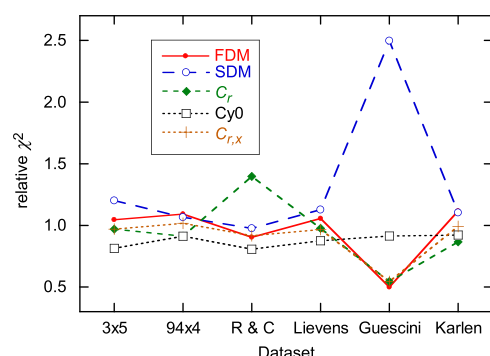


Fig. 15. χ^2 values from the weighted calibration fits for the different C_q markers, normalized to unity for each dataset. In each case, common weights were used for the 5 $C_{q,s}$.

cycles, for which C_q calibration is relevant.

In C_q calibration fitting, weighting is important anytime the data include dilutions where Poisson scatter is significant, and may be needed when other sources of anomalous scatter can be identified in the C_q data. Individual reaction growth curves generally show much more scatter in the plateau region than in the baseline cycles, again indicating the need for weighting when fitting to whole-curve models. However, we find that such weighting makes little difference in the quality of the C_q estimates, as evidenced from their subsequent calibration fitting. This means that the dispersion in C_q estimates is dominated by factors other than data noise. We have suggested that for concentrations where Poisson scatter is negligible, pipette volume uncertainty may be the main such factor [12,13,15] – a notion that has also received attention in more recent work by de Ronde et al. [36]. The good news: qPCR is capable of much better precision with experimental procedures that reduce such operational uncertainties. A cautionary note is in order here. We have attributed the variance discrepancy to excess ensemble variance, which assumes that the fitted data are truly raw. There are indications that some qPCR data have been smoothed before being reported. Such smoothing can lead to falsely precise parametric SEs, making the latter the source of the variance discrepancy [12,34]. There is no easy way to get reliable statistics by fitting smoothed data, so qPCR instrument makers should ensure that the raw data are accessible, and users should use these data in any fitting.

Which C_q marker is "best"? Using χ^2 from the calibration fitting as the goodness metric, the best performer varies with dataset, as shown in Fig. 15, where the values for each dataset have been normalized to the average for that dataset. SDM is somewhat poorer than the others, even when the anomalously high χ^2 for the Guescini data is excluded. Cy0 is best, with mean = 0.87 and low dispersion, $\sigma = 0.05$. (However, the low σ is also a statistical quirk, as without the high SDM value, Cy0 would be high for the Guescini data.) $C_{r,x}$ is close behind Cy0, at 0.90(18); and C_r and FDM are in statistical agreement – 0.94(27) and 0.95(23). In short, with quality routines for estimating them, there is little difference in performance among the compared C_q markers, with slightly poorer results for SDM.

The results shown in Fig. 15 are based on actual performance on the replicate datasets. A possibly more telling answer to the "best C_q " question is given by the SE comparisons in Fig. 14, because these represent what *could* be achieved without extraneous experimental sources of variability. As we have noted earlier, all but $C_{r,x}$ are affected by limitations in the whole-curve model, as manifested in systematic trends in the LS fit residuals. $C_{r,x}$ is much less sensitive to such effects, because only the first few growth cycles are included in the fitting. This restricted cycle range also means that weights will probably never be needed in these fits, while they are often warranted in whole-curve fitting and will make a difference there when data noise is the only source of variance. It is for these reasons that we chose $C_{r,x}$ in our method of estimating N_0 from C_q variance [15].

Competing interests

The authors declare that there are no conflicts of interest.

Acknowledgments

We thank Bob Rutledge for providing the raw data for the 15 reactions from [7] – the 3×5 data used here. Funding has been provided to ANS by grant Sp721/4-2 of the Deutsche Forschungsgemeinschaft (DFG).

Appendix A. Supplementary data

Supplementary data associated with this article can be found, in the online version, at <https://doi.org/10.1016/j.bdq.2019.100084>.

References

- [1] D. Svec, A. Tichopad, V. Novosadova, M.W. Pfaffl, M. Kubista, How good is a PCR efficiency estimate: recommendations for precise and robust qPCR efficiency estimates, *Biomol. Detect. Quantif.* 3 (2015) 9–16.
- [2] M.W. Pfaffl, A new mathematical model for relative quantification in real-time RT-PCR, *Nucleic Acids Res.* 29 (2001) e45.
- [3] J. Tellinghuisen, Using nonlinear least squares to assess relative expression and its uncertainty in real-time qPCR studies, *Anal. Biochem.* 496 (2016) 1–3.
- [4] S.A. Bustin, Absolute quantification of mRNA using real-time reverse transcription polymerase chain reaction assays, *J. Mol. Endocrinol.* 25 (2000) 169–193.
- [5] R.G. Rutledge, C. Cote, Mathematics of quantitative kinetic PCR and the application of standard curves, *Nucleic Acids Res.* 31 (2003) e93.
- [6] A. Larionov, A. Krause, W. Miller, A standard curve based method for relative real time PCR data processing, *BMC Bioinformatics* 6 (2005) 62.
- [7] R. Rutledge, D. Stewart, Critical evaluation of methods used to determine amplification efficiency refutes the exponential character of real-time PCR, *BMC Mol. Biol.* 9 (2008) 96.
- [8] J. Tellinghuisen, Simple algorithms for nonlinear calibration by the classical and standard additions method, *Analyst* 130 (2005) 370–378.
- [9] J. Tellinghuisen, Least squares in calibration: weights, nonlinearity, and other nuisances, *Method Enzymol.* 454 (2009) 259–285.
- [10] J.M. Ruijter, M.W. Pfaffl, S. Zhao, A.N. Spiess, G. Boggy, J. Blom, R.G. Rutledge, D. Sisti, A. Lievens, K. De Preter, S. Derveaux, J. Hellemans, J. Vandesompele, Evaluation of qPCR curve analysis methods for reliable biomarker discovery: Bias, resolution, precision, and implications, *Methods* 59 (2013) 32–46.
- [11] J. Tellinghuisen, A.-N. Spiess, Statistical uncertainty and its propagation in the analysis of quantitative polymerase chain reaction data: comparison of methods, *Anal. Biochem.* 464 (2014) 94–102.
- [12] J. Tellinghuisen, A.-N. Spiess, Bias and imprecision in analysis of real-time quantitative polymerase chain reaction data, *Anal. Chem.* 87 (2015) 8925–8931.
- [13] J. Tellinghuisen, A.-N. Spiess, Comparing real-time quantitative polymerase chain reaction analysis methods for precision, linearity, and accuracy of estimating amplification efficiency, *Anal. Biochem.* 449 (2014) 76–82.
- [14] A.-N. Spiess, S. Rödiger, M. Burdukiewicz, T. Volksdorf, J. Tellinghuisen, System-specific periodicity in quantitative real-time polymerase chain reaction data questions threshold-based quantitation, *Sci. Rep.* 6 (2016) 38951.
- [15] J. Tellinghuisen, A.-N. Spiess, Absolute copy number from the statistics of the quantification cycle in replicate quantitative polymerase chain reaction experiments, *Anal. Chem.* 87 (2015) 1889–1895.
- [16] G.J. Boggy, P.J. Woolf, A mechanistic model of PCR for accurate quantification of quantitative PCR data, *PLoS One* 5 (2010) e12355.
- [17] A.C. Carr, S.D. Moore, Robust quantification of polymerase chain reactions using global fitting, *PLoS One* 7 (2012) e37640.
- [18] I. Chervoneva, Y.Y. Li, B. Iglewicz, S. Waldman, T. Hyslop, Relative quantification based on logistic models for individual polymerase chain reactions, *Stat. Med.* 26 (2007) 5596–5611.
- [19] R. Rutledge, D. Stewart, A kinetic-based sigmoidal model for the polymerase chain reaction and its application to high-capacity quantitative real-time PCR, *BMC Biotechnol.* 8 (2008) 47.
- [20] A.N. Spiess, C. Feig, C. Ritz, Highly accurate sigmoidal fitting of real-time PCR data by introducing a parameter for asymmetry, *BMC Bioinformatics* 9 (2008) 221.
- [21] J. Tellinghuisen, Can you trust the parametric standard errors in nonlinear least squares? Yes, with provisos, *Biochim. Biophys. Acta* 1862 (2018) 886–894.
- [22] J. Tellinghuisen, Statistical error propagation, *J. Phys. Chem. A* 105 (2001) 3917–3921.
- [23] M. Guescini, D. Sisti, M.B. Rocchi, L. Stocchi, V. Stocchi, A new real-time PCR method to overcome significant quantitative inaccuracy due to slight amplification inhibition, *BMC Bioinformatics* 9 (2008) 326.
- [24] A. Tichopad, M. Dilger, G. Schwarz, M.W. Pfaffl, Standardized determination of real-time PCR efficiency from a single reaction set-up, *Nucleic Acids Res.* 31 (2003) e122.
- [25] P.R. Bevington, *Data Reduction and Error Analysis for the Physical Sciences*, McGraw-Hill, New York, 1969.

- [26] J. Tellinghuisen, A Monte Carlo study of precision, bias, and non-Gaussian distributions in nonlinear least squares, *J. Phys. Chem. A* 104 (2000) 2834–2844.
- [27] J.D. Ingle Jr., Exchange of comments: signal flicker noise and noise power spectra, *Anal. Chem.* 49 (1977) 339–340 (and references therein).
- [28] Q.C. Zeng, E. Zhang, H. Dong, J. Tellinghuisen, Weighted least squares in calibration: estimating data variance functions in high-performance liquid chromatography, *J. Chromatogr. A* 1206 (2008) 147–152.
- [29] J. Tellinghuisen, C.H. Bolster, Least-squares analysis of phosphorus soil sorption data with weighting from variance function estimation: a statistical case for the Freundlich isotherm, *Environ. Sci. Technol.* 44 (2010) 5029–5034.
- [30] J. Tellinghuisen, Weighted least-squares in calibration: what difference does it make? *Analyst* 132 (2007) 536–543.
- [31] S. Zhao, R.D. Fernald, Comprehensive algorithm for quantitative real-time polymerase chain reaction, *J. Comput. Biol.* 12 (2005) 1047–1064.
- [32] A. Lievens, S. Van Aelst, M. Van den Bulcke, E. Goetghebeur, Enhanced analysis of real-time PCR data by using a variable efficiency model: FPK-PCR, *Nucleic Acids Res.* 40 (2012) e10.
- [33] Y. Karlen, A. McNair, S. Perseguers, C. Mazza, N. Mermod, Statistical significance of quantitative PCR, *BMC Bioinformatics* 8 (2007) 131.
- [34] A.-N. Spiess, C. Deutschmann, M. Burdukiewicz, R. Himmelreich, K. Klat, P. Schierack, S. Rödiger, Impact of smoothing on parameter estimation in quantitative DNA amplification experiments, *Clin. Chem.* 61 (2015) 379–388.
- [35] M. Burdukiewicz, A.-N. Spiess, K.A. Blagodatskikh, W. Lehmann, P. Schierack, S. Rödiger, Algorithms for automated detection of hook effect-bearing amplification curves, *Biomol. Detect. Quantif.* 16 (2018) 1–4.
- [36] M.W.J. de Ronde, J.M. Ruijter, D. Lanfear, A. Bayes-Genes, M.G.M. Kok, E.E. Creemers, Y.M. Pinto, S.-J. Pinto-Sietsma, Practical data handling pipeline improves performance of qPCR-based circulating miRNA measurements, *RNA* 23 (2017) 811–821.