

Weighted Linear Regression based Data Analytics for Decision Making after Early Failures

Ross, Robert; Ypma, P.A.C.; Koopmans, Gerben

DOI

[10.1109/ISGTAsia49270.2021.9715569](https://doi.org/10.1109/ISGTAsia49270.2021.9715569)

Publication date

2021

Document Version

Final published version

Published in

IEEE PES Innovative Smart Grid Technologies Asia (IEEE PES ISGT-Asia 2021)

Citation (APA)

Ross, R., Ypma, P. A. C., & Koopmans, G. (2021). Weighted Linear Regression based Data Analytics for Decision Making after Early Failures. In *IEEE PES Innovative Smart Grid Technologies Asia (IEEE PES ISGT-Asia 2021)* <https://doi.org/10.1109/ISGTAsia49270.2021.9715569>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Green Open Access added to TU Delft Institutional Repository

'You share, we take care!' - Taverne project

<https://www.openaccess.nl/en/you-share-we-take-care>

Otherwise as indicated in the copyright section: the publisher is the copyright holder of this work and the author uses the Dutch legislation to make this work public.

Weighted Linear Regression based Data Analytics for Decision Making after Early Failures

Robert Ross^{1,2}

¹ IWO

Ede, Netherlands

² Faculty EWI

TU Delft

Delft, Netherlands

ORCID 0000-0003-3803-0953

Peter Ypma^{1,2}

¹ Faculty EWI

TU Delft

Delft, Netherlands

² IWO

Ede, Netherlands

P.Ypma@tudelft.nl

Gerben Koopmans

IWO

Ede, Netherlands

G.Koopmans@iwo.nl

Abstract—This paper discusses a method for data analysis of early failures that are typically due to defects from production and/or installation. The approach is based on Weighted Linear Regression and aims at estimating the next failure moment and the failure probability in an interval. As an additional goal, methods are intended to be suitable for light computing devices.

Keywords—early failures, Weibull, bathtub curve, censored data, hazard rate, weighted linear regression, similarity index, asymptotic behavior, emergency

I. INTRODUCTION EARLY FAILURE DECISION-MAKING

The present paper discusses situations where decision-making requires data analytics on relatively small and/or incomplete data sets such as with (very) early failures. An experienced typical example is a power grid connection that is supposed to last for > 40 years, but multiple cable joints fail in the first year. It is noteworthy that components were factory tested and that the connection passed a full commissioning test program. If the cause is not obvious from forensics (as yet), fears may grow that many more failures may follow.

Resilience is the capability to recover from mishap like these unforeseen early failures. The present work supports grid resilience by developing a Weibull data analyzer for light computing devices to provide statistical grounds for decision-making. It aims at estimating three often relevant quantities: (a) the expected moment of a next failure; (b) the probability of failures in a time interval; (c) the number of future early failures. The first two are subject of the present paper.

Having sufficient data supports accurate analysis, but speed may be preferred over accuracy if decisions about repair or massive replacement are urgent. As an additional objective in our project, the methodology is aimed to be implementable in light computing devices, in order to be widely accessible to utility work force, to small and medium-sized businesses (SMEs) and for possible implementation into firmware of smart devices.

The paper firstly discusses competing processes, bathtub curves and the effect of fast aging due to defective (however as yet functioning) components. This leads to a preliminary conclusion that early failures can be treated by a fairly simple distribution, but with a possibly unknown sample size.

After a brief description of asymptotic behavior, censored data, ranking indices and similarity index, the progress with

respect to weighted linear regression (WLR) is discussed. A model is presented that accurately approximates the LR-weights by a set of asymptotic power functions and is suitable for light computing. Alternative approaches for estimating next failure times, interval probabilities and total sample size n are explored in parallel. The results will be compared to the present findings in later publications.

II. COMPETING MECHANISMS AND BATHTUB CURVES

The failure behavior of products in terms of the hazard rate $h(t)$ is often described with a bathtub curve. Such curves are generally stated to consist of a competition of teething (also called ‘child mortality’), random failure and wear out failure. Assuming failure processes can be described with Weibull distributions, the hazard rate of each process is:

$$h(t) = \frac{\beta \cdot t^{\beta-1}}{\alpha^\beta} \quad (1)$$

Here α and β are the scale and shape parameter of the applicable Weibull distribution. The teething, random failure and wear-out processes are commonly associated with shape parameter $\beta < 1$, $\beta = 1$ and $\beta > 1$ respectively. The total hazard rate for a normal product batch h_n is the sum of the competing processes: h_t , h_r and h_w respectively:

$$h_n = h_t + h_r + h_w \quad (2)$$

The sum of these hazard rates yields the well-known bathtub curve, cf. Fig. 1. The bathtub curve represents the three typical, competing processes of teething, random failure and wear-out. Competition of failure mechanisms is also found in systems that are built from multiple components that each can make the system fail. For instance, a cable circuit will at least consist of switchgear, cable termination, cable, cable joints and more. If any of these fails, the circuit fails.

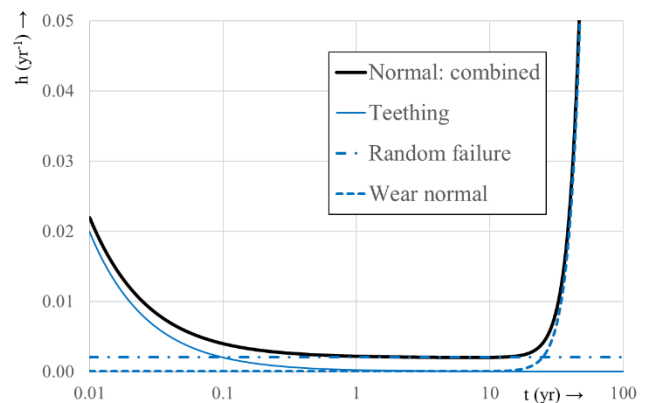


Fig. 1. Typical bathtub curve for competition between teething, random failure and wear-out failure. The three hazard rates accumulate.

The research on data analysis of early failures is funded by the Netherlands Ministry of Economic Affairs and Climate through the RvO agency, Grant ref. nr. TEUE418008, TKI Project FINDGO.

The research for approximation of weights for linear regression received support from the EU H2020 R&I program and RvO under ECSEL grant agreement No 826417 (Project Power2Power).

Quality control reduces teething by pre-aging products before delivery. This is usually done at enhanced stresses to accelerate aging (particularly of the teething phenomenon if possible) and get beyond the moment where teething dominates. For the case in Fig. 1, pre-aging might aim for the equivalent of $t=0.1$ yr where teething has decayed to the same level as random aging. The surviving products then form a high reliability batch (Ch.14, Fig. 4 in [1]).

III. EARLY FAILURES AND MIXED BATCHES CASES

If early failures do occur despite quality control and commissioning tests, these faults are usually not due to teething, but rather to fast aging imperfections or defects. A second bathtub with a faster wear-out appears. Teething and random failure are assumed to be the same for simplicity.

Fig. 2 compares the two bathtub curves. As mentioned above, the difference is in the wear-out: hazard rate h_n depends on h_w in the end, cf. (2), but on h_{fw} for the fast-wearing defect products. The total hazard rate h_d of the defect subgroup is:

$$h_d = h_t + h_r + h_{fw} \quad (3)$$

A flawed batch of products typically contains two subgroups: a fraction p_n that consists of normal and a fraction $p_d=1-p_n$ that consists of defective products. The ratio of these fractions is denoted as N:D = p_n/p_d . The subgroups have a reliability function R_n respectively R_d . The combined hazard rate h_c in this case is (cf. Section 2.5.2 in [2]):

$$h_c = \frac{p_n \cdot h_n \cdot R_n + p_d \cdot h_d \cdot R_d}{p_n \cdot R_n + p_d \cdot R_d} \quad (4)$$

If both p_n and p_d are non-zero, the combined hazard rate h_c forms a double bathtub (see Fig. 3). In the present example, the early failures become noticeable after about $t = 0.2$ yr. The wear-out process Weibull shape parameter $\beta > 1$, but the scale parameter α is much smaller than normal as a result of the imperfections. The defect subgroup will become depleted at some moment and normal products will remain (unless $p_n=0$). The combined hazard rate h_c will then drop to the level of random failure of the surviving (normal) products.

The assumption of a specific distribution (here: Weibull) is a big advantage in the reliability analysis. Still, the statistical analysis of (4) is complicated due to the involved 8 Weibull parameters (for the 4 processes) plus the ratio N:D. As in Fig. 3, strongly deviant behavior is often dominated by a single process, which simplifies the analysis. As an important consequence, early failures can often be attributed to a single temporarily dominant process. As a consolation, early failures can often be described with a single Weibull distribution.

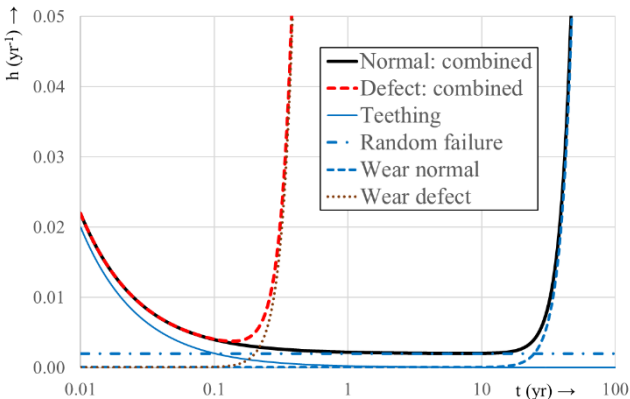


Fig. 2. Bathtub curves for a batch of only normal or defect products.

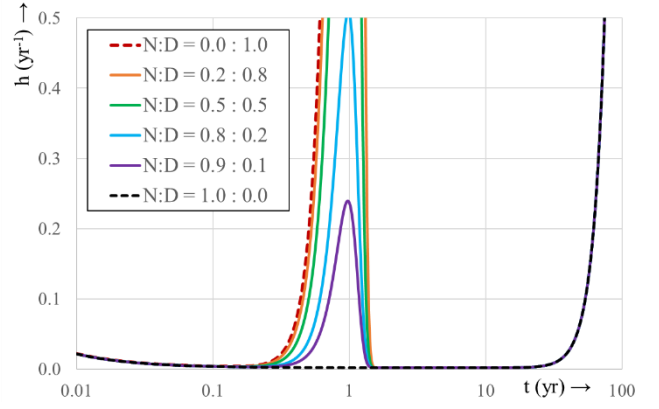


Fig. 3. Bathtub curves for various fractions of defect products. N:D stands for the ratio p_n/p_d . With mixed batches of normal and defective products a double bathtub appears. Quality control normally removes the part before $t = 0.1$ yr. The defect failure behavior then becomes the first dominant effect.

Moreover, if random failures occur in parallel, they are often caused by external impact and may easily be recognized as such (e.g. a cable failing due to digging). In Fig. 3, the failures that occur shortly after about $t = 0.2$ yr, are practically only due to the fast wear-out process with hazard rate h_{fw} . As mentioned above: $\beta_{fw} > 1$, but α_{fw} (much) smaller than normal.

If forensic analyses and/or diagnostics are not conclusive (as yet), the challenge is to analyze failure data for the sake of grid resilience. The target is to do this by light computing.

IV. BACKGROUND OF APPLIED METHODS

Now, three separate subjects are described that are instrumental to developing the data analysis as discussed in the present paper: ranking adjustment with censored data sets, asymptotic behavior and the similarity index.

A. Censored Data and Non-Integer Indices

When (very) early failures occur suspicion arises that the delivered batch is flawed. To what extent (i.e. N:D) may remain unknown. Data analytics is applied to answer basic questions about the next failure time and the probability of new failure in a given interval.

At the moment of evaluation, it is expected – or at least held possible – that more failures are to follow. The data analytics should therefore reckon with censored data, i.e. times of failures that are as yet unknown. If all components have been commissioned at once and kept in operation equally long, this is right censoring. If the components differ in operational time, then the censoring is likely random.

With random censoring, the ranking of the observed failure times may be disturbed. This is the case if some components have been shorter in operation than already observed failure times. Both in plotting and linear regression the ranking is often adjusted [3]. The observed failure times are then attributed with adjusted ranking indices $I(i)$ that are not necessarily integers. The procedure is [3]:

$$I(i) = I(i-1) + \frac{n+1-I(i-1)}{n+2-C_i} \quad (5)$$

Here i is the original ranking index of failure times t_i as observed; $I(i)$ the adjusted ranking index, C_i the original index i increased with the number of censored failure times $< t_i$. By definition $I(0)=0$. This adjustment is based on the number of permutations of ranking and adopted by IEEE and IEC [3].

Alternative ranking adjustments exist and/or are under development. E.g., intervals between observed failure times may be considered (e.g., [4]), using the idea that the chance to find an as yet censored time in that interval is proportional to the difference in F -value (F being the cumulative probability). However, at this stage the IEC standard [3] is followed.

B. Asymptotic Behavior as a Power Function

One criterion for good estimators is consistency. This implies that the estimated parameter E_n asymptotically approaches a true value E_∞ with increasing n . Such asymptotic behavior can often be described in terms of a power function model with parameters P , Q and R (not to be confused with probability or a reliability function), cf. Section 6.1.3 in [2]:

$$E_n = E_\infty + Q \cdot (n - R)^P \quad (6)$$

E_n asymptotically approaches E_∞ if the power P is negative and a singularity appears at $n = R$ consequently (unless $Q = 0$). This asymptotic model was used for the regression weights.

C. Similarity Index

A Similarity Index S_{fg} is used to quantify how similar two (normalized) distribution densities f and g are, (cf. Section 3.6 in [2]). If f and g are identical, $S_{fg}=1$, if f and g have nothing in common $S_{fg} = 0$. The definition of S_{fg} is:

$$S_{fg} = \frac{\langle f \cdot g \rangle}{\langle f \cdot f \rangle + \langle g \cdot g \rangle - \langle f \cdot g \rangle} \quad (7)$$

Here $\langle f \cdot g \rangle$ is an inner product of two probability densities or mass functions that is usually defined as a time integral or an indexed summation. S_{fg} can also be used over intervals for comparison. E.g., the similarity of observations and specifications can be evaluated for an interval up to τ_e by $S_{fg}(0, \tau_e)$ and compared to $S_{fg}(0, \infty)$ extrapolating the trend to infinitely [5]. The significance of S_{fg} increases with lifetime [6]. In the presently discussed work, S_{fg} is used for testing the distribution of the model regression weights against theory.

V. WLR PARAMETER ESTIMATION

The most widely applied parameter estimator families are Linear Regression (LR) and Maximum Likelihood (ML). The focus in this paper is on Weighted Linear Regression (WLR). An advantage of (W)LR is that the estimated parameters, the errors and the confidence limits are analytically obtained. This is particularly useful for light computing. Secondly, the least squares (LS) estimates and the best fit in a plot are fully consistent. These advantages are not fundamental or absolute arguments, but do serve the goal of light computing.

A. Weighted Linear Regression

(W)LR is based on a linear relationship between plotting position Z and $\ln(t)$ from which α and β are estimated:

$$Z(p) = \ln(-\ln(1 - p)) = \beta \cdot \ln(t) - \beta \cdot \ln \alpha \quad (8)$$

Here p is a probability, i.e., a value of the Weibull cumulative distribution $F(t; \alpha, \beta)$. If preferred, the first 'ln' in each term can be replaced by $^{10}\log$, which differs by a factor $^{10}\log(e)$. Parameters are estimated by an LS method. The terms LR and LS commonly refer to the same parameter estimators. Ordinary LS (OLS) minimizes the sum of deviations of actual data and fit. WLR or WLS is somewhat more accurate than OLS by attributing weights w_i to each observation.

In a Weibull plot, the observed $\ln(t_i)$ are plotted against the expected value $\langle Z_i \rangle$. Contrary to most practice in graphs,

probability plots have variables $\langle Z_i \rangle$ scaled along the vertical axis and covariables $\ln(t_i)$ along the horizontal axis. WLR estimators a_{WLR} and b_{WLR} of α respectively β are:

$$a_{WLR} = \exp\left(\overline{\ln(t)}_w - \frac{\overline{\langle Z \rangle}_w}{b_{WLR}}\right) \quad (9)$$

$$b_{WLR} = \frac{\overline{(\langle Z \rangle - \overline{\langle Z \rangle}_w)^2}_w}{\overline{((\langle Z \rangle - \overline{\langle Z \rangle}_w) \cdot (\ln t - \overline{\ln t}_w))}_w} \quad (10)$$

The suffix w means that the averages are weighted, i.e., the weighted average $\bar{u}_w$ of a series of observations u_i ($i=1, \dots, n$) is:

$$\bar{u}_w = \frac{\sum_{i=1}^n (w_i \cdot u_i)}{\sum_{i=1}^n w_i} \quad (11)$$

With OLS, for all i : $w_{i,n}=1$. With WLR, the $w_{i,n}$ are taken as the inverse variances v_i of the plotting positions Z_i :

$$w_{i,n} = \frac{1}{v_{i,n}} = \frac{1}{\langle Z_{i,n}^2 \rangle - \langle Z_{i,n} \rangle^2} \quad (12)$$

The smaller variance $v_{i,n}$, the heavier weighs observation t_i .

B. Calculation of LR Weights

For given n , the variance $v_{i,n} = \text{var}(Z_{i,n})$ can be calculated as a summation (Equation 2.7 in [7]). For larger i and n the calculation becomes demanding and such a summation is not suitable for $\text{var}(Z_{i,n})$ with (often non-integer) adjusted ranking indices I . A look-up table and round I to an integer can be used [3]. However, $\langle Z_i(p) \rangle$ can also be assessed in the p -domain, cf. (8), and the Beta distribution $B(p)$. This allows non-integer parameters x, y and indices I . The Beta density function f_B is:

$$f_B(p; x, y) = \frac{\Gamma(x+y)}{\Gamma(x) \cdot \Gamma(y)} \cdot p^{x-1} \cdot (1-p)^{y-1} \quad (13)$$

The expected j^{th} moment of $Z_{I,n}$ is:

$$\langle Z_{I,n}^j \rangle = \int_0^1 Z^j(p) \cdot f_B(p; I, n+1-I) dp \quad (14)$$

In terms of which the variance $v_{I,n}$ is given by:

$$v_{I,n} = \langle Z_{I,n}^2 \rangle - \langle Z_{I,n} \rangle^2 \quad (15)$$

Z is singular at $p = 0$ and 1 , implying that care is required in numerical integrations if I is close to 1 or n . As for $I = 1$, the summation shows that $v_{1,n} = \pi^2/6$ for every n .

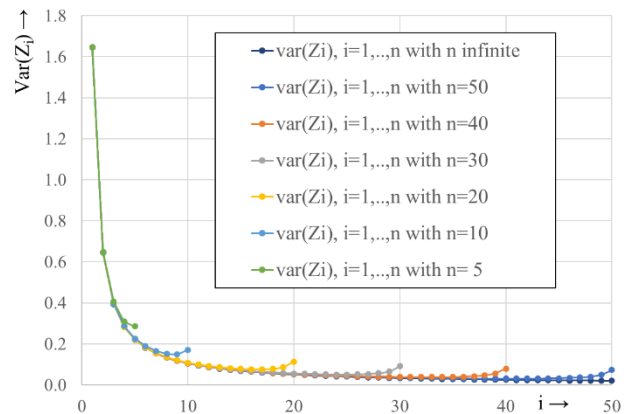


Fig. 4. For various sample sizes $n=5, \dots, 50$ and $n \rightarrow \infty$, the variances $\text{var}(Z_i)$ with $i \leq n$. The inverse variances are the weights for regression analysis.

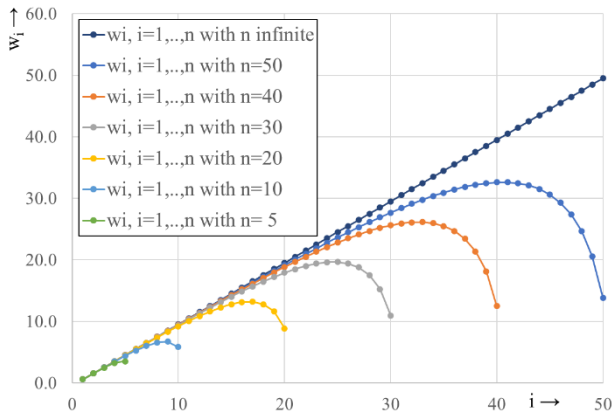


Fig. 5. For various sample sizes $n=5, \dots, 50$ and $n \rightarrow \infty$, the weights w_i with $i \leq n$. The small upward curl in the v_i curve causes a significant drop in w_i .

The domain was subdivided into k segments concentrated near the endpoints, where Z is singular. The segment boundaries were chosen at $[\cos(i\pi/k) + 1]/2$ with $i = 0, \dots, k$. Within the segments 64-point Gauss-Legendre quadrature was used, in which the integral is approximated as a weighted sum of integrand values (see pp. 887 and 918 of [8]). The variances v_i were calculated for all $i=1, \dots, n$ with various n in the range $[1, 2000]$ and used as references for a model to fit w_i .

Fig. 4 and Fig. 5 show the distribution of v_i respectively w_i for sample sizes $n=5, 10(10)50$ and $n \rightarrow \infty$. ‘ $A(B)C$ ’ means a sequence from A to C with an increment B . For finite $n > 6$, the v_i distributions curl up at both ends (i.e., at $i \downarrow 1$ and $i \rightarrow n$). As a consequence, the w_i distributions appear to have a maximum for finite $n > 6$ in the range $(n/2) \leq i \leq n$ (cf. Fig. 5).

C. The Multiple Asymptotic Model

In order to obtain a model that is suitable for both light computing and for non-integer indices I , the asymptotic behavior was studied in detail. After an explorative research on for $n \leq 200$ yielded two parametric models [9], the merits of power functions for approximating the variances and weights were investigated in greater depth, up to $n = 2000$.

Analyzing the v_i distribution, it was firstly observed that for $n \rightarrow \infty$ and finite i , the variances v_i can be described as a power series:

$$v_{i,n \rightarrow \infty} = \frac{\pi^2}{6} - \sum_{j=2}^i \frac{1}{(j-1)^2} \quad (16)$$

Describing the asymptotic behavior with a power function, for large (but finite) i , the $v_{i,n \rightarrow \infty}$ can be largely described with a power function with $(P, Q, R) = (-1, 1, 0.5)$. For small i , adding a second power function with $(P, Q, R) = (-3, 0.1, 0.3445)$ yields a very good approximation of (16). The $w_{i,n \rightarrow \infty}$ appear almost linear with finite i (Fig. 5). The $v_{i,n \rightarrow \infty}$ are approximated:

$$v_{i,n \rightarrow \infty} \approx (i - 0.5)^{-1} - 0.1 \cdot (i - 0.3445)^{-3} \quad (17)$$

The weights significantly drop for finite n and $i \rightarrow n$. A large variety of asymptotic relations n , i and $n-i$ were explored for $v_{i,n}$ and $w_{i,n}$ yielding power functions with varying success. A very successful model for $v_{i,n}$ with finite n appeared an extension of (17) with mixed power functions of n , $n-i$ and i :

$$v_{i,n} \approx (i - .5)^{-1} - 0.1 \cdot (i - 0.3445)^{-3} + (0.125 \cdot (n + 0.343)^{-1.656}) \cdot (n + 0.8 - i)^{-0.75} \cdot (i - 1)^{1.4} \quad (18)$$

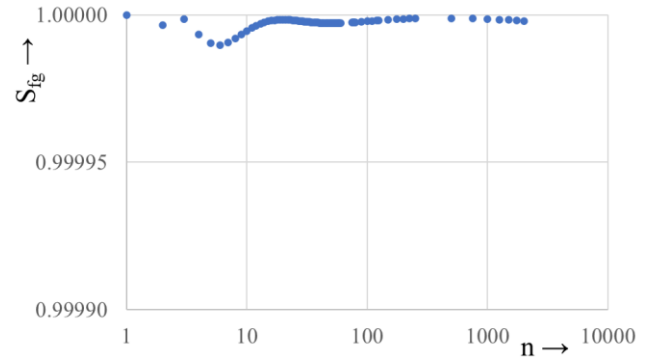


Fig. 6. Similarity between theoretical weights and the (18) model.

The $w_{i,n}$ by this model were tested against the theoretical $w_{i,n}$ for $n=1(1)60, 75(1)80, 80(10)120, 125(25)250(250)2000$. A total of 27750 $w_{i,n}$ -values were involved in the test.

For all investigated $n \leq 2000$, the S_{fg} appeared > 0.9999885 which is very close to 1. The largest deviation between the model and theory was found for $n = 6$, where $1 - S_{fg} = 1.02 \cdot 10^{-5}$ (see Fig. 6). For each individual $w_{i,n}$, the absolute error is $< 1\%$ for $n \leq 500$; and $< 2.7\%$ for $n \leq 2000$. The model is suitable for light computing as well as for non-integer ranking index I .

With these weights, the Weibull parameters for the best WLR fit follow from (9) and (10).

VI. TIME AND PROBABILITY TO NEXT FAILURE

Within the present approach, the Weibull distribution with the estimated parameters is the starting point for inferences. Other approaches are being explored as well, such as directly estimating expected quantities like next failure time $\langle t_{r+1} \rangle$ from the data. Which is to be preferred and why, is subject of the program and debate. Here, the distribution is determined by: (a) observed and censored data; (b) adjusted ranking I through (5); (c) the weights w_i through (18); (d) estimated a_{WLR} and b_{WLR} through (9)-(11) (summing over w_i of known t_i).

As for censoring, if components are installed by replacement or in various non-simultaneous projects, the start of operation and lifetimes will vary. Failure of components will then lead to random censoring. Equation (5) can be used for adjusting rankings and estimating the ruling Weibull distribution. It depends on the individual cases how complicated the forecast of failures will be.

In the following, however, the focus is on cases where all products were put into service simultaneously such as joints in a cable circuit. This is called right censoring, i.e. r out of n objects have failed with $r < n$. The censored failure times t_i ($r < i \leq n$) are larger than a minimum time θ . So, we assume that the moment of evaluation is sometime after the last failure t_r . So, as for the next failure time t_{r+1} :

$$t_{r+1} > \theta \geq t_r \quad (19)$$

The expected next failure time $\tau \equiv \langle t_{r+1} \rangle$ and confidence limits $t_{i,A\%}$ are studied for two cases: without and with taking the (19) condition into account. It is also noted that the inverse Weibull distribution relates time t and probability F :

$$t(F) = \alpha \cdot [-\ln(1 - F)]^{\frac{1}{\beta}} \quad (20)$$

So, θ is associated with a F_θ . It is noteworthy, that in this approach the times do depend on n , which may be unknown. If so, reasonable assumptions must be made about this n .

A. Next failure time from the Beta distribution

Without taking (19) into account, the expected next failure time τ follows from the Beta density (cf. (13)), (20) and known or estimated Weibull parameters α and β as:

$$\tau_{i=r+1} = \int_0^1 t(p) \cdot f_B(p; r+1, n-r) dp \quad (21)$$

The $A\%$ confidence limits $t_{r+1,A\%}$ are found by firstly determining $F_{A\%}$. The Beta distribution $B(F_{A\%}; x, y)$ is:

$$B(F_{A\%}; x, y) = \int_0^{F_{A\%}} f_B(p; x, y) dp = A\% \quad (22)$$

The inverse Beta distribution $B^{inv}(A\%; x, y)$ yields $F_{A\%}$ for given $A\%$ and is commonly available in spreadsheets. The $A\%$ limits $t_{r+1,A\%}$ are found with known or estimated n as:

$$t_{i=r+1,A\%} = t(B^{inv}(A\%; r+1; n-r)) \quad (23)$$

In a Weibull plot, $\ln(t)$ is used and $\ln(t_{r+1})$ can be determined in a similar fashion as in (21). The $\Delta A\%$ Beta confidence intervals between two $A\%$ confidence limits (i.e. the quantiles) define a range where the $r+1^{\text{th}}$ failure is expected with $\Delta A\%$ probability. These intervals are generally very wide due to ignoring the already observed failure times.

B. Next failure time conditional on previous failures

Taking the (19) condition and relation (20) into account, the next $n-r$ failures will have associated probabilities p with: $F_\theta < p < 1$. Unordered future failures are uniformly distributed over $[F_\theta, 1]$. The ranked failures are Beta distributed over this range with coordinate q :

$$q = \frac{(p - F_\theta)}{(1 - F_\theta)} \quad (24)$$

Let the index of ranked future failures be $j = 1, \dots, n-r$. Failure $i = r+1$, i.e. $j = 1$, is the first out of $n-r$, as yet censored, failures. The applicable Beta density is $f_B(q; 1, n-r)$. With the condition $t_j > \theta$, the expected next failure time $\tau_\theta \equiv \langle t_{j=1} \rangle$ and the $A\%$ limits $t_{j=1,A\%,\theta}$ are found from (20)-(23) after elaborating p in terms of q using (24) and substituting $f_B(q; 1, n-r)$ for $f_B(p; r+1, n-r)$. Consequently:

$$\tau_\theta = \int_0^1 t(F_\theta + q \cdot (1 - F_\theta)) \cdot f_B(q; 1, n-r) dq \quad (25)$$

$$t_{j=1,A\%,\theta} = t(F_\theta + (1 - F_\theta) \cdot B^{inv}(A\%; 1; n-r)) \quad (26)$$

VII. DISCUSSION AND CONCLUSIONS

This study is initiated after some high impact incidents in the Dutch power grid. This urged to review the concepts of teething and early failures. E.g., 5 joints failed after 58, 78, 90, 100 and 107 days (cf. Section 9.4.3 in [2]). The decision to replace >100 joints was proven right by forensics afterwards.

The commonly encountered bathtub curve for the hazard rate of product batches implies that multiple processes are active. This means that the total failure distribution is a mix of multiple distributions as well. A proper set of factory tests and commissioning tests should reduce the teething problems, after which a well performing batch should remain.

If early failures occur despite factory and commissioning tests, these may be due to a subgroup of defective components

that wear abnormally fast. If a Weibull distribution applies to these faults, then probably $\beta > 1$ and α is much smaller than specified for this process. Typically, the components seem to work well for a short period and then the faults occur with decreasing intervals (as in the case of the five early joint faults mentioned above). Though the overall statistics seem complicated, early failures can often be treated as a single distribution as its hazard rate dominates other processes. This makes the statistical analysis relatively simple. However, the total sample size of the subgroup of defects may be unknown.

A situation of early failures often calls for a decision to preventively replace or continue to repair. Either choice can have a big impact, whether it is the right choice or not. This applies to large utilities and small and medium enterprises (SMEs) alike. The challenge is taken up to develop methods that are accessible for a wide audience, which is translated into developing methods for light computing devices.

One objective was to develop WLR with an accurate weight approximation that is suitable for light computing and for non-integer ranking indices. With (18), this is achieved for sample sizes up to $n = 2000$. The similarity of two distributions can be tested with the similarity index, S_{fg} (7). S_{fg} was also used to compare the distributions of the observed and approximated weights. These proved to be very similar. The WLR is built into a spreadsheet that is published as freeware for educational and non-commercial use [10].

As discussed here, the approach is to first estimate the ruling distribution for a single dominant failure mechanism and conduct inferences based on that. After the WLR estimation of α and β , the methods aim at estimating the time-to-next failure and confidence intervals. Typical expected times and limits follow from (21)-(23), but these are very wide and do not acknowledge already observed failure times. Taking the known previous failure time t_r into account, leads to alternative estimates for the expected next time and the confidence limits through (25)-(26).

REFERENCES

- [1] R. Ross and G. Koopmans, "Reliability and Degradation of Power Electronic Materials," in Reliability of Organic Compounds in Microelectronics and Optoelectronics, Cham, Springer Nature Switzerland AG, 2021.
- [2] R. Ross, Reliability Analysis for Asset Management of Electric Power Grids, Hoboken, NJ: Wiley-IEEE Press, 2019.
- [3] IEC TC112, "IEC 62539(E):2007 Guide for the statistical analysis of," International Electrotechnical Committee, Geneva, 2007.
- [4] W. Wang, "Refined rank regression method with censors," *Quality and Reliability Engineering International*, vol. 20, nr. 7, pp. 667-678, November 2004.
- [5] R. Ross, "Evaluating Methods for Detecting End-of-Life and Non-Compliance of Asset Populations," in *Cigré SCD1 Colloquium*, Rio de Janeiro, 2015.
- [6] P. A. Ypma and R. Ross, "Determining the Similarity between Observed and Expected Ageing Behavior," in *Proceedings ICEMPE*, Xi'an, 2017.
- [7] J. S. White, "The Moments of Log-Weibull Order Statistics," *Technometrics*, vol. 11, no. 2, pp. 373-386, May 1969.
- [8] M. a. S. I. A. National Bureau of Standards (Abramowitz, Handbook of Mathematical Functions with Formulas, Graphs, and Mathematical Tables, New York: Dover Publications, 1972.
- [9] P. A. Ypma and R. Ross, "Approximation of weight numbers for Weibull Least Squares Parameter," in *Proceedings QR2MSE 2019*, Zhangjiajie, 2019.
- [10] IWO, "IWO Data Analyzer Freeware," IWO, November 2021. [Online]. Available: www.iwo.nl.