

Improving CubeSat reliability: Subsystem redundancy or improved testing?

Bouwmeester, J.; Menicucci, A.; Gill, E.K.A.

DOI

[10.1016/j.ress.2021.108288](https://doi.org/10.1016/j.ress.2021.108288)

Publication date

2022

Document Version

Final published version

Published in

Reliability Engineering & System Safety

Citation (APA)

Bouwmeester, J., Menicucci, A., & Gill, E. K. A. (2022). Improving CubeSat reliability: Subsystem redundancy or improved testing? *Reliability Engineering & System Safety*, 220, Article 108288. <https://doi.org/10.1016/j.ress.2021.108288>

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.



Improving CubeSat reliability: Subsystem redundancy or improved testing?

J. Bouwmeester^{*}, A. Menicucci, E.K.A. Gill

Delft University of Technology, Kluyverweg 1, 2629HS, Delft, The Netherlands

ARTICLE INFO

Keywords:

CubeSat reliability
Subsystem redundancy
Bayesian inference
Reliability models

ABSTRACT

The objective of this paper is to investigate which approach would lead to more reliable CubeSats: full subsystem redundancy or improved testing. Based on data from surveys, the reliability of satellites and subsystems is estimated using a Kaplan–Meier estimator. Subsequently, a variety of reliability models is defined and their maximum likelihood estimates are compared. A product of a Lognormal distribution addressing immaturity failure and a Gompertz distribution addressing wear-out is found to best represent CubeSat reliability. Bayesian inference is used to find realistic wear-out parameters by using failure data of small satellites. Subsystem reliability estimates are subsequently found using a similar approach. A reliability model for CubeSats with redundant subsystems is established, verified and applied in a Monte Carlo simulation. The results are compared with a model for reduced immaturity failure. Allocating resources to reduction of immaturity failures through improved testing is considered to be superior to allocating these resources to the implementation of subsystem redundancy.

1. Introduction

This paper investigates the reliability of CubeSats and a key choice for improving this. The CubeSat concept was introduced in 1999 by Puig-Suari and Twiggs [1] of the California Polytechnic State University. A CubeSat is a satellite with standardized mechanical interfaces for a launch interface adaptor and comprises one or more units of $10 \times 10 \times 10$ cm (1U). First a literature study will be provided on satellite reliability statistics, redundancy in CubeSats and development approaches to increase reliability. From this, the research question will be provided and motivated.

1.1. Satellite reliability data and statistics

According to a study of 156 satellite failures of all mass classes [2], 41% occurred in the first year of operations. Another study, focusing on small satellite reliability which analyses the satellite failure and anomalies of subsystems of 222 satellites up to 500 kg, provides reliability over operational lifetime [3]. One of the conclusions is that satellites below 10 kg show a relatively high infant mortality rate and short lifetime compared with satellites between 10 kg and 500 kg. Telemetry, tracking and command (TT&C), the Thermal Control System (TCS), and the mechanisms and structures (M&S) contribute most to infant mortality, while the EPS contributes to the largest number of failures overall [3]. In a statistical study on the first 100 launched CubeSats performed in 2013 [4], mission failures of CubeSats are analysed on

a high level. One major conclusion is that for a third of all failed missions radio signals have never been received after launch. In a study of Swartwout in 2017, focusing on university class satellites [5], it is found that a quarter of these missions are dead-on-arrival. The universities that produce multiple spacecraft show significant improvement in success. Swartwout states, based on personal experience, that “student-led projects often fail because of lack of time/resources given to systems-level testing” [5]. A NASA report presents statistical analysis on CubeSat reliability [6] based on the database from Swartwout [7]. Out of the 390 CubeSats analysed, only the data of 21 CubeSats was deemed appropriate and complete for use in statistical regression. A Weibull fit is provided, but it is also concluded that the lack of useful data makes practical reliability estimation difficult [6]. According to a recent extensive study in 2019 on 855 CubeSats by Villela et al. [8], mission success rates have risen from 30% in 2005 and levelled off to approximately 75% in 2018. Infant mortality is found to be dominant overall. The associated public database [9], however, does not provide dates of failure or dates related to the listed operational status. It is therefore not suitable for reliability modelling. Langer established a CubeSat failure database and performed statistical analysis in 2016 [10]. This database partially originates from a cooperative survey with the first author of this paper. It was first used in a study on CubeSat electrical interfaces [11] and is further extended by Langer through literature search and individual correspondence. Public available databases from Swartwout [7] and Kulu [12] have been used

^{*} Corresponding author.

E-mail address: jasper.bouwmeester@tudelft.nl (J. Bouwmeester).

<https://doi.org/10.1016/j.ress.2021.108288>

Received 2 June 2021; Received in revised form 9 December 2021; Accepted 19 December 2021

Available online 8 January 2022

0951-8320/© 2022 The Authors. Published by Elsevier Ltd. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

to address CubeSat teams and to validate some data. The study presents a maximum likelihood estimate of CubeSats and their subsystems using a Weibull mixture and a 'percentage-non-zero' component to address dead-on-arrival specifically [10]. The associated (non-public) database provides the most comprehensive failure data on CubeSats available for further study, comprising data on 178 satellites after sanitation of all input.

1.2. Redundancy to improve reliability

A common approach to reliability in larger spacecraft is to apply redundancy to critical subsystems. When redundancy is applied to the full subsystem unit(s) and its components (e.g. sun sensors), it is possible to use cross-strapping between them using redundant electrical interfaces. This is for example applied on the Flying Laptop Platform (FLP) satellite [13] and the Mars Reconnaissance Orbiter [14]. This is a small satellite of 117 kg and a medium sized satellite of 2180 kg respectively. Both have sufficient volume to implement such a strategy. The limited volume of CubeSats typically prohibits such an advanced redundancy concept. Delfi-C³, a 3U CubeSat of Delft University of Technology, offered redundancy in the form of a non-identical backup system for the acquisition and transmission of data of one its demonstration payloads [15]. This system was however not fully worked out during launch due to lack of time. On its successor Delfi-n3Xt, a single-point-of-failure free design philosophy was initially applied to the critical subsystems for the key mission objectives [16]. This yields redundancy of the on-board computer, the radio transceivers and parts of the electrical power subsystem. Other subsystems had partial redundancy, such as a simplified backup system for attitude determination and control to secure de-tumbling. Especially the non-identical redundancy was complex and time-consuming to implement. Also, full single-point-of-failure free design was not achieved, while the time to test out all the arbitration systems was limited. The experiences with these satellites on the implementation of redundancy is a key motivation to investigate if this is the right strategy to improve reliability.

1.3. Development strategies to improve reliability

Many papers on CubeSat reliability focus on lessons learnt from development and flight experience. Menchinelli et al. provides an approach to improve CubeSat reliability with the aid of Failure Mode, Effects and Criticality Analysis (FMECA) using a detailed step-wise approach [17]. Pantoji et al. compared several failure identification tools for their CubeSat: Fault Tree Analysis (FTA), Failure Mode Effect Analysis (FMEA), FMECA, Probability Risk Assessment (PRA) and Risk Response Matrix (RRM) [18]. They considered RRM most useful for their project because of its simplicity which was a main criterion for a student-driven project. Nieto-Peroy & Emami used the experience of several university class CubeSats to address lessons learnt in terms of team, procurement, procedures, testing and cooperation to be able to reduce associated failure probability [19]. Venturini et al. [20] investigated CubeSat reliability using interviews of CubeSat developers in the United States. Out of the 27 anomalies discussed during the interviews, it was expected that 19 of them could have been avoided with more ground testing. Based on the shared experiences, development guidelines are formulated with a strong emphasis on (more extensive) testing [20]. Furthermore, a warning is provided that commercial-off-the-shelf subsystems and components are not always compliant to the specifications. Doyle et al. present findings of a survey with eight responses from different CubeSat developers, indicating that mission level testing is done by a majority of the teams but limited to hours or weeks [21]. It is concluded that guidelines for comprehensive system level testing, including worst case and failure scenarios, is desired. Berthoud et al. analysed the project management of university class CubeSats through a few CubeSat case studies [22]. It provides lessons learnt on a proper organization with a leading staff and motivated students to increase mission success and emphasizes the need for extensive testing.

1.4. Motivation and research question

According to literature, reliability of CubeSats is a concern and failures often occur early in life; in the first months while the mission design life time is a year or more. Applying redundancy is a commonly applied measure to improve system reliability in larger spacecraft. Experience with CubeSats is limited and indicates extensive resources are needed to implement redundancy properly. CubeSats have limited volume and financial budgets. Earlier studies on CubeSat innovations propose lean electrical interfaces [23,24] and integration of main satellite bus subsystems into one physical unit for additional payload volume [25]. Such an architectural strategy would conflict with subsystem redundancy. Financial, technical and human resources, which would be needed for implementing redundancy, can alternatively be allocated to improving the reliability of the individual subsystems and interfaces through more extensive testing. The suggestions in literature are to improve development practices, with an emphasis on system level testing. These qualitative studies however do not investigate the quantitative impact of improved testing. Hard conclusions of which approach is better, considering limited resources, can therefore not yet be made. This leads to the following research question for this study:

What leads to higher CubeSat reliability over its mission life time: full subsystem redundancy or improved testing?

Additionally, a combination of these measures and potential further improvement through iterative satellite development will also be investigated.

2. Scope and data

In this section, the scope of the study, failure classification and the available data is explained.

2.1. Scope

The scope of this reliability study is limited to critical failures of CubeSats caused by failing subsystems including their physical interfaces. This means that the failure within a subsystem causes loss of the satellite or its mission [26]. CubeSats are considered to have a limited mission scope for which typically all subsystems are essential to function properly. If a critical subsystem is not able to recover from its failure, only redundancy can help to mitigate satellite failure. In order to model the reliability of a satellite within the defined scope and investigate the effect of subsystem redundancy and improved testing, a model for irrecoverable subsystem failures is needed.

2.2. Failure classification

For this study, a classification is desired which divides failures into a limited set of groups which have specific behaviour in terms of failure rate over time and can be linked to the two satellite failure reduction strategies, subsystem redundancy and improved testing (see Section 1.4). The bathtub curve is a widely used theoretical reliability model for the failure rate of a group of devices. Its origin is unknown, but it is described many times for hundreds of years [27]. An example of the bathtub curve is shown in Fig. 1.

The bathtub curve starts at the roll-out of operations of a system with a declining failure rate. Failures due to poor design, production errors of components, too limited testing and/or wrong analysis of the operational environment may lead to early failure in life. For this reason, these failures are often called 'infant mortality'. Subsequently follows a period of random failures, represented by a constant failure rate. At a later lifetime of the system, wear-out failures are shown with increasing failure rate. Eventually every satellite subsystem will wear-out due to accumulation of environmental effects and/or internal ageing throughout its active operation. Wear-out is thus inevitable on the long term irrespective of the maturity of the development.

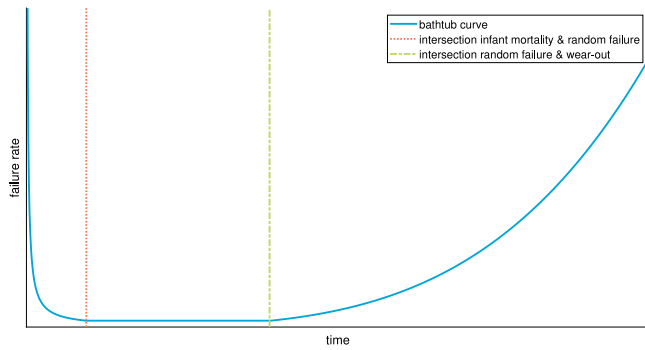


Fig. 1. Schematic bathtub failure rate curve.

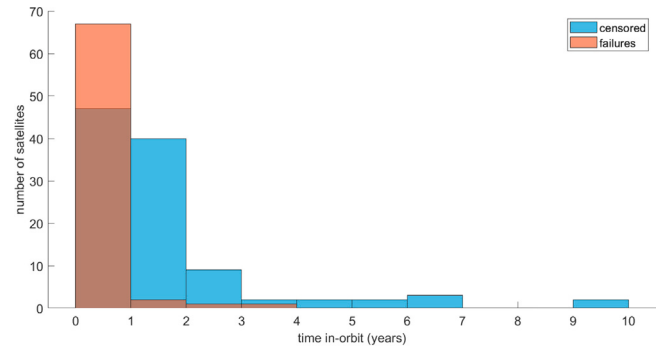


Fig. 2. Histogram of data from CubeSat failure database.

Although the bathtub curve in Fig. 1 provides a continuous model over time, it has a few issues. When using adjacent time windows, each class of failures is cut-off in time. This is not realistic for any of the failure classes. Moreover, the required boundary conditions, to avoid discontinuities in failure rate, will complicate the estimation of model parameters. A compound of continuous models over time, instead of adjacent, is therefore preferred. In this case, however, the tails of the decreasing and increasing failure rate models may be difficult to distinguish from a constant failure rate model. Furthermore, according to Klutke et al. there is a lack of empirical evidence for the theoretical bathtub curve [27] and according to Wong 'there is no such thing as random failure' [28]. If a specific random event occurs frequently, e.g. many times per day, the majority of satellites are developed such that they are able to survive this. However, some may fail. If a specific random event occurs early in life, the term 'infant mortality' applies. If the average time between such random events increases, the potential failure rate will decrease and the period in which a failure can be expected increases. This makes the term 'infant mortality' less suitable. There are many different events which may cause failure and the time between these events varies. The events may be environmental (e.g. particle radiation), deterministic (e.g. a software state) or user-imposed (e.g. change of operation by telecommand). They have in common that they can in principle be mitigated by extensive analysis, on-ground and in-orbit testing followed by necessary improvements. A widely accepted term for a long-term decreasing failure rate does not exist and new terms will be subject to debate. It is however helpful for this paper to continue with a brief term. For this paper, the term 'immaturity failure' is introduced which comprises infant mortality as well as long-term failures which exhibit a decreasing failure rate. For this study, immaturity failures and wear-out failures together constitute all possible satellite failures.

2.3. Selected failure data

For reliability modelling in this study, CubeSat failure and survival data is needed at specified times in orbit. As described in Section 1.1, the 'CubeSat failure database' from Langer [10] fulfils this objective and is most comprehensive after sanitation compared to other candidates [7,12]. The database contains 71 observed failures from 178 CubeSats for an observation window from the 20th of May 2003 to the 31st of December 2014. It contains censor times and failure times at satellite and subsystem level. A failure indicates the loss of satellite operability or loss of a main mission objective. A censor time indicates that the satellite was turned off intentionally after achieving mission success, de-orbited after successful operations or was still operational at the time of inquiry. A histogram of the data is shown in Fig. 2. While a satellite failure is relatively easy to identify by the operators and provides robust data, allocation to a specific subsystem is less trivial. The full CubeSat Failure Database cannot be disclosed due

to agreements with some interviewed persons, but a significant part is provided in a public database [29]. A few examples, with their subsystem allocation between brackets, are:

- *The determined mission failure cause is a polarity of magnetic sensor became inverted. [attitude determination and control]*
- *The determined satellite failure cause is degradation of solar cells, which finally lead to negative energy budget. [electrical power sub-system]*
- *For the satellite failure, we have one or more hypotheses: on-board boot code was overwritten during a hardware reset or software crash or both [command and data handling]*
- *The determined failure cause is unknown [unspecified]*

Key hypotheses are used for allocation of failures to subsystems, which introduces some sensitivity to the subjective interpretation of these observations. This is considered acceptable, provided that the subsystem model estimation does not depend very strongly on this limited set of subsystem failure observation. The sampling method provided in Section 5.3 accounts for this.

The initial survey database [29] also contains some information on mission design life times. They range from one month to five years for launched CubeSats. Langer et al. already provided reliability model estimates for CubeSats and their main subsystems [10]. While these models are good fits and are useful for high level analysis, they were unfit for the purpose of this study as will become clear in Section 3.3.

A second database containing small satellite failures data, already used in another study at TU Delft by Guo and Monas [3], is also available for this study. It contains data on 152 satellites launched between 1990 and 2010 with a mass lower than 500 kg and reports 83 failures and 69 censored items.

The small satellite database will be used to check whether the model selection for CubeSat failures can be generalized to other classes of satellites. Furthermore, the fact that in this database there are relatively more failures beyond one year compared to the CubeSat failure database, as can be seen in Fig. 3, will be used as input for Bayesian inference as explained in Section 4.2.

3. Method

Based on the defined failure classification, several steps are required to answer the research question. These steps are provided in Fig. 4. This method results in a simulation based on empirical CubeSat reliability data.

3.1. Non-parametric reliability model

The first step of the reliability analysis is to censor the data and apply the Kaplan–Meier Estimator (KME) [30]. This is a non-parametric

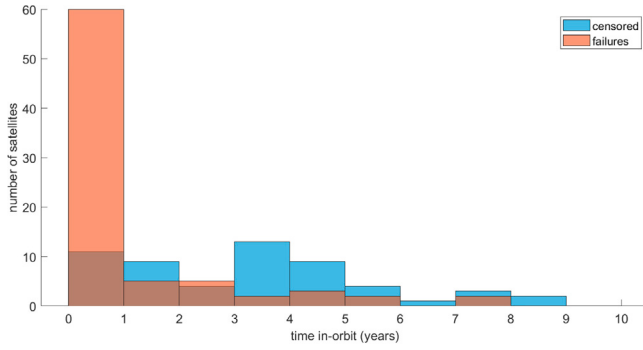


Fig. 3. Histogram of data from small satellite failure database.

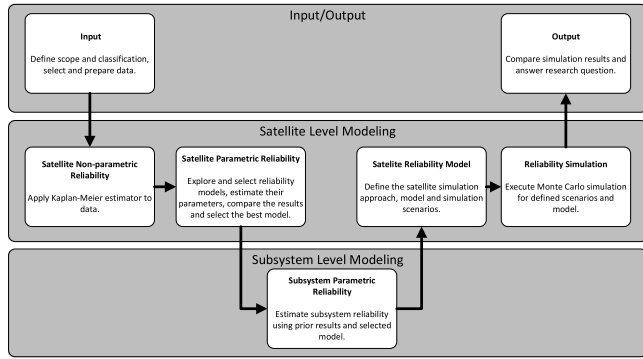


Fig. 4. Flow diagram of reliability analysis.

survival function ($\hat{S}$) which can be directly established from failure and censor times of satellite and subsystem data [31,32].

$$\hat{S}(t) = \prod_{i: t_i < t} \left(1 - \frac{d_i}{n_i}\right) \quad (1)$$

KME provides the estimated survival function over time t and is updated at each time t_i that a number of failures d_i occur, having n_i operational units at risk. The variance of the estimator can be calculated by Greenwood's formula [33]. The confidence interval (CI) can subsequently be determined by applying the α -quantile of the normal distribution $z_{\alpha/2}$. For a 95% confidence interval, $z_{\alpha/2} = z_{0.025} = 1.96$.

$$\text{var}(\hat{S}(t)) = \hat{S}(t)^2 \sum_{i: t_i < t} \left(\frac{d_i}{n_i(n_i - d_i)} \right) \quad (2)$$

$$CI = \hat{S}(t) \pm z_{\alpha/2} \sqrt{\text{var}(\hat{S}(t))} \quad (3)$$

KME is useful as first step as it can act as input for a least squares estimator and as a reference to assess the goodness-of-fit of a parametric model.

3.2. Parametric reliability models

In contrast to non-parametric models, parametric models provide a smooth distribution over time to 'fit' the empirical data. The basic metric is the reliability as function of time $R(t)$, similar to $\hat{S}(t)$ of KME, representing the expected fraction of survivors over time. From this, other measures of reliability can be derived such as the probability density $f(t)$ and the failure rate $\lambda(t)$ [34].

$$f(t) = -\frac{dR(t)}{dt} \quad (4)$$

$$\lambda(t) = \frac{f(t)}{R(t)} \quad (5)$$

Table 1

Overview of basic reliability models.

Weibull:	gamma:
$R(t) = \exp \left[-\left(\frac{t}{\theta} \right)^\beta \right]$	$R(t) = 1 - \frac{\int_0^t t^{\beta-1} e^{-t} dt}{\int_0^\infty t^{\beta-1} e^{-t} dt}$
Gompertz:	log-logistic:
$R(t) = \exp \left[-\eta \left(e^{\left(\frac{t}{\theta} \right)} - 1 \right) \right]$	$R(t) = 1 - \frac{1}{1 + (t/\theta)^\beta}$
log-normal:	
$R(t) = 1 - \frac{1}{\sigma\sqrt{2\pi}} \int_0^t \frac{2}{x} \exp \left[-\frac{1}{2} \frac{(\ln(x) - \mu)^2}{2\sigma^2} \right] dx$	

Previous research on satellite lifetime reliability has focused primarily on the Weibull distribution [3,10,35]. As indicated in Section 2, there is an interest in models capable of representing both immaturity failure and wear-out. This investigation is complemented with the gamma, Gompertz, log-logistic and log-normal distributions (see Table 1) which are often applied in survival analysis in general [36,37] to identify if they would yield better distributions than the Weibull. The gamma, log-normal and log-logistic distributions have a long right tail and are therefore considered less suitable for wear-out. The Weibull distribution [38] can be used for 'immaturity failures' when $\beta < 1$. When $\beta > 2$ the probability density starts concave upward, which is considered a suitable boundary condition for wear-out phenomena. For the gamma distribution, the failure rate is monotonically decreasing with $\beta \leq 1$. The Gompertz distribution [39] can only be used to address wear-out failures. When $\eta \geq 0.1$, the probability density at $t = 0$ is already significant which would interfere with immaturity failure so a boundary condition of $\eta \leq 0.1$ is deemed appropriate. The log-logistic distribution can be used for immaturity failures with a monotonically decreasing failure rate $\lambda(t)$ with shape factor $\beta < 1$. For the log-normal distribution, the mode can be skewed towards $t = 0$ when $\sigma^2 \gg \mu$. This can be used to model immaturity failure.

To model immaturity failure and wear-out, a compound of the basic reliability models is needed. A mixture of n Weibull distributions is a common method applied in previous satellite reliability studies [10,35] using α_i as the normalized weight factor for each component i :

$$R(t) = \sum_{i=1}^n \alpha_i R_i(t) \quad (6)$$

A mixture model divides all devices in populations, where the weight factor α can be regarded as the probability of failing according to a basic model, for example immaturity failure or wear-out. Castet & Saleh have applied the Weibull-Weibull mixture in their study [35]. They use the maximum error and average error over time between the Weibull distribution and the non-parametric data as the benchmark and proved a good quality of the fit. However, the estimate also yields a reliability of still 0.86 after 100 years which is considered unrealistic. For this study, long term behaviour is considered important. A potential problem with mixtures is that wear-out is not effective on the population allocated to immaturity failure, which may create a problem if the latter is not diminished in time. Therefore, next to mixtures of different basic models, also an alternative compound is investigated. Such alternative to a mixture, is a product of the basic reliability models. The basic reliability models would then act as reliability components put in series and means that all devices are subject to the risk of both immaturity failure and wear-out:

$$R(t) = \prod_{i=1}^n R_i(t) \quad (7)$$

A reliability product requires that the immaturity failure component leaves sufficient survivors before the peak of the probability density of the wear-out component to avoid that the latter is superfluous. This is a key difference with respect to reliability mixtures, which requires the opposite.

Reliability mixtures and reliability products are both considered candidates for modelling satellite reliability assuming both immaturity failure and wear-out. With four basic models for immaturity failure, two for wear-out and two types of compounds (mixtures and products), a total of 16 combinations are under investigation.

3.3. Satellite model estimation

There are various ways to estimate the parameters of a probability function based on a set of empirical data, such as least squares, maximum likelihood and Bayesian inference.

For varying parameters vector θ failure times t_i , the least squares estimate θ_{LSE} can be calculated by minimizing the sum of the residuals SS_{res} :

$$SS_{res} = \sum_{i=1}^n \left(\hat{S}(t_i) - R(t_i|\theta) \right)^2 \quad (8)$$

$$\theta_{LSE} = \arg \min SS_{res}(\theta|t). \quad (9)$$

When also using also the censor times t_j , the maximum likelihood estimate θ_{LSE} can be calculated by maximizing the likelihood L :

$$L(\theta|t) = \prod_{i=1}^n f(t_i|\theta) \cdot \prod_{j=1}^m R(t_j|\theta) \quad (10)$$

$$\theta_{MLE} = \arg \max L(\theta|t) \quad (11)$$

For some reliability models, the probability density at $t = 0$ is zero. When failures are present at exactly $t = 0$, so will the likelihood and as a consequence the MLE will fail. The CubeSat and small satellite failure databases have a few of those entries. These satellites were either dead-on-arrival or have worked up to a few hours but failed before the first possible ground station contact. To mitigate the MLE issue and to account for limited potential operational lifetime, a bias of +0.1 day is applied for reported events at $t = 0$. Although this value is arbitrary, it is sufficiently high to avoid computational issues while it is insignificantly small compared to the data set observation window of many years.

Bayesian inference is based on the Bayes theorem [40], where a prior belief in the form of a probability distribution of the model parameters $P(\theta)$ is introduced to calculate the posterior distribution of those parameters:

$$P(\theta|t) = \frac{L(t|\theta) \cdot P(\theta)}{\int L(t|\theta) \cdot P(\theta) d\theta} \propto L(t|\theta) \cdot P(\theta). \quad (12)$$

The integral in the denominator, which ensures integration of all values to one, is constant and can be ignored if an improper posterior is allowed. Similar to the MLE, the Maximum-a-Posteriori (MAP) can be calculated.

$$\theta_{MAP} = \arg \max P(\theta|t) \quad (13)$$

For censored satellite reliability data, MLE is found to be a suitable method for parametric estimation [10,32,41]. Bayesian inference can even lead to improved estimation [3]. However, it requires prior information on the model as a uniform prior will otherwise lead to the same result as the MLE. First an appropriate model needs to be found, for which MLE is chosen as estimator. If it is possible to establish an informed prior, Bayesian will subsequently be applied.

3.4. Comparing model quality

To compare the quality of the different models, several criteria are used.

Criterion 1: Goodness-of-fit

The goodness-of-fit can be determined by the coefficient of determination R^2 based on the sum of the squares of the residuals SS_{res} divided by the sum of squares w.r.t. the mean SS_{tot} . An R^2 of 1 indicates a perfect fit and a value of 0 indicates an uncorrelated fit [42]. The adjusted R^2 is used [43] which includes a penalty for the number of parameters k in relation to the number of observations n .

$$R^2 = 1 - \frac{SS_{res}}{SS_{tot}} = 1 - \frac{\sum_{i=1}^n \left(\hat{S}(t) - R(t_i|\theta) \right)^2}{\sum_{i=1}^n \left(\hat{S}(t) - \bar{\hat{S}} \right)^2} \quad (14)$$

$$R_{adj}^2 = 1 - \left[\frac{(1 - R^2)(n - 1)}{n - k - 1} \right] \quad (15)$$

The vast majority of the maximum likelihood estimates of the single basic models as provided in Section 4.1 yield $R_{adj}^2 \geq 0.95$. A compound model should in principle have a similar or better goodness of fit, so $R_{adj}^2 \geq 0.95$ is used as acceptance criterion for selecting an appropriate compound model.

Criterion 2: Mode of wear-out component

For the wear-out the mode $Mode_2$ should not occur too early or unrealistically late. For this study a range between 1 and 25 years is chosen as acceptance criterion.

Criterion 3: Long-term reliability

A good compound model assures that the vast majority of the satellites have worn out in a reasonable time. The acceptance criterion used in this study is that the reliability $R_{t=50y} \leq 0.1\%$.

Criterion 4: Wear-out shape parameter boundary

For both the CubeSat and small satellite failure databases there are many failures in the first year, which causes the MLE to naturally converge towards immaturity failure values for both components. To ensure that the second component of the compound model addresses wear-out, a boundary condition $B.C.$ is used in the MLE for Weibull ($\beta_2 \geq 2$) and Gompertz ($\eta_2 \leq 0.1$). Ideally the MLE does not converge to this boundary condition but finds a true (local) maximum.

Criterion 5: Akaike information criterion

The best ranking criterion for models using MLE is the likelihood L . However, likelihood can be improved by adding parameters with the risk that they provide little extra information. To deal with this issue, the Akaike Information Criterion (AIC) [44] can be used where the number of parameters k is included as a penalty. The AIC value has no meaning in absolute sense, but the AIC of different models can be compared in relative sense where a lower value is better:

$$AIC = 2k - 2 \cdot \ln(L). \quad (16)$$

Taking the best model as reference, all models which have an AIC value of +6 or higher are rejected as this yields a likelihood ratio of $< 5\%$ compared to the best model in terms of AIC based on an equal number of parameters.

The model with the lowest AIC, which meets the previously described acceptance criteria, is chosen as the preferred model. The focus of this study is the CubeSat failure database. However, the small satellite failure database is also analysed.

4. Satellite reliability model estimates and selection

The next step is to perform model estimation and selection to provide a realistic CubeSat reliability failure distribution.

Table 2
Results of the MLE estimates.

CubeSat failure database					
Model	AIC	R^2_{adj}	Mode ₂	$R_{t=50y}$	B.C.
Lognormal single	−95.4	0.959		34.4%	
Logn.-Gomp. mix.	−93.7	0.904		1.36%	Yes
Logn.-Weib. mix.	−92.3	0.914	56	34.2%	No
Logn.-Gomp. prod.	−91.4	0.958	287	34.4%	No
Logn.-Weib. prod.	−91.4	0.958	50	14.4%	No
Loglogistic single	−89.1	0.962		33.1%	
Logl.-Gomp. mix.	−87.6	0.901	14	2.45%	Yes
Logl.-Weib. mix.	−86.1	0.919	21	6.19%	No
Weibull single	−86.1	0.953		27.2%	
Logl.-Weib. prod.	−85.1	0.961	20	0.00%	No
Logl.-Gomp. prod.	−85.1	0.961	39	0.00%	No
Weib.-Gomp. mix.	−83.9	0.944	18	0.12%	Yes
Weib.-Weib. mix.	−83.3	0.949	90	44.3%	No
Gamma single	−82.3	0.913		17.4%	
Weib.-Weib. prod.	−82.1	0.952	115	27.2%	No
Weib.-Gomp. prod.	−82.1	0.952	37	0.00%	No
Gam.-Weib. prod.	−78.3	0.911	30	0.00%	No
Gam.-Gomp. prod.	−78.3	0.911	300	17.4%	No
Gam.-Gomp. mix.	−78.1	0.955	60	43.6%	No
Gam.-Weib. mix.	−78.1	0.955	52	24.2%	No
Small satellite failure database					
Model	AIC	R^2_{adj}	Mo ₂	$R_{t=50y}$	B.C.
Logn.-Gomp. mix.	171.5	0.985	17	3.97%	Yes
Logn.-Weib. prod.	171.6	0.984	18	0.25%	No
Logn.-Gomp. prod.	172.1	0.982	21	0.00%	No
Lognormal single	172.6	0.949		26.7%	
Logl.-Weib. mix.	174.1	0.987	12	4.32%	Yes
Logl.-Weib. prod.	176.1	0.981	19	0.11%	No
Logn.-Weib. mix.	176.1	0.970	5.3	24.3%	No
Loglogistic single	176.3	0.951		25.9%	
Weib.-Gomp. mix.	176.4	0.982	17	0.03%	Yes
Logl.-Gomp. prod.	176.5	0.979	21	0.00%	No
Weibull single	177.4	0.959		19.9%	
Weib.-Gomp. prod.	179.8	0.972	22	0.00%	No
Gamma single	180.3	0.954		12.5%	
Logl.-Gomp. mix.	180.6	0.967	6.1	23.7%	No
Weib.-Weib. prod.	181.4	0.958	332	19.9%	No
Gam.-Weib. mix.	181.7	0.972	12	0.01%	Yes
Gam.-Gomp. mix.	182.6	0.963	16	0.00%	Yes
Weib.-Weib. mix.	182.8	0.963	5.4	19.3%	No
Gam.-Weib. prod.	183.8	0.955	24	0.00%	No
Gam.-Gomp. prod.	184.0	0.955	24	0.00%	No

4.1. Results & comparison of model estimates using MLE

The results of the chosen set of compound models (Section 3.2) for both databases, using the maximum-likelihood-estimator, are provided in Table 2. The list is ordered on AIC value, with the lowest (i.e. best) at the top. If the acceptance criteria for the values are violated, the specific value is coloured red.

The best AIC values for both databases are provided by the Lognormal distribution for immaturity failure. Figs. 5 and 6 provides the reliability curves for those. Another conclusion which can be made is that MLE results for mixtures have a tendency to converge to the boundary condition for the wear-out shape factor and they yield unrealistically high reliability on the long term ($R_{t=50y} > 1\%$). In this respect, reliability products are performing significantly better. For the CubeSat failure database, which is the main focus of this study, unfortunately all models violate at least one of the acceptance criteria. The single Lognormal model even has the best AIC and also one of the highest R^2_{adj} . This indicates that the CubeSat failure database does not contain sufficient information (relatively low number of failures beyond the first year) to provide convincing results of the wear-out parameter values. For the small satellite failure database however, the Lognormal-Gompertz product meets all acceptance criteria and comes very close to the highest AIC. It combines a good fit for immaturity failure with the

Table 3
Results of the MAP estimates for Cubesat failure database.

Estimate	AIC	R^2_{adj}	Mode ₂	$R_{t=50y}$
CubeSat MLE	−91.4	0.958	287.1	34.35%
75% of w	−91.2	0.950	21.1	0.00%
w = 20/152	−91.1	0.951	21.3	0.00%
125% of w	−91.1	0.946	21.3	0.00%
Small sat. MAP	−85.0	0.889	21.2	0.00%

property of a relatively short survivor tail for the wear-out compared to other models.

4.2. Results for model estimates using Bayesian inference

As there is insufficient data on CubeSat failures to provide appropriate estimates using the MLE method, another approach is investigated. Using Bayesian inference, the small satellite failure database could be used to provide a prior distribution for CubeSats based on the fact that CubeSats are a subset of the class of small satellites. The model parameters for CubeSats should lie within a range of values which have a reasonable likelihood ratio compared to the maximum likelihood estimate of the small satellites. First, the posterior distribution of the small satellite database is calculated using a uniform prior. With this uninformed prior, the maximum-a-posteriori (MAP) parameters should match the maximum likelihood estimate (MLE). It also provides a four-dimensional posterior (because of having four parameters) which can be converted to a prior for the CubeSat failure database. The normalized marginal posterior distribution, which is the sum for one parameter over the other parameter values normalized to its maximum, are provided in Fig. 7. It should be noted that the MAP is found in a four-dimensional posterior, and may therefore differ from the maximum of the marginal distribution.

It can be seen that the MAP values in Fig. 7 closely match the MLE values in Fig. 6 as expected, with very minor differences which are fully explained by the grid step size for the Bayesian inference.

Using the posterior distribution of the small satellite failure database directly as prior for the CubeSat failure database would effectively yield the same result as combining both databases into one. There are only 20 CubeSats in the small satellite database of 152 satellites, significantly less than the 178 in the CubeSat database. The posterior should therefore first be weakened to be able to act as prior. There is no common method or clear rules as Bayesian inference relies on the additional knowledge and/or logical reasoning to define the prior. The posterior is proportional to the likelihood times the prior as explained in Eq. (12) and the likelihood is the product of probabilities of failure or survival. Given this product relationship, a weakening of each value of the posterior using an weight exponent w between 0 and 1 would be a natural choice to achieve a new prior:

$$p_{prior}(\theta) = p_{post}(\theta)^w. \quad (17)$$

Table 3 and Fig. 9 provide the results in comparison with the MLE for the CubeSat and the MAP of the small satellites. The estimate using $w = 20/152$ meets all acceptance criteria. Its lognormal parameters come very close to the MLE, which is expected due the relative high number of early failures. The wear-out parameters are now realistic w.r.t. the MLE, while the AIC comes very close. The overall distribution does not show a high sensitivity to the chosen translation weight when changed by $\pm 25\%$. Fig. 10 shows the marginal distributions of the CubeSat posterior, the used prior and the posterior of the small satellites which the prior is based upon (see Fig. 8).

In conclusion, the best posterior distribution for the CubeSat failures is found through Bayesian inference on the CubeSat failure database using the small satellite failure database posterior as input which is translated to a prior using Eq. (17) with $w = 20/152$. The resulting maximum-a-posteriori estimate is a Lognormal-Gompertz product with parameters $\mu_1 = 1.35$, $\sigma_1 = 6.30$, $\theta_2 = 4.7$ and $\eta_2 = 0.0107$.

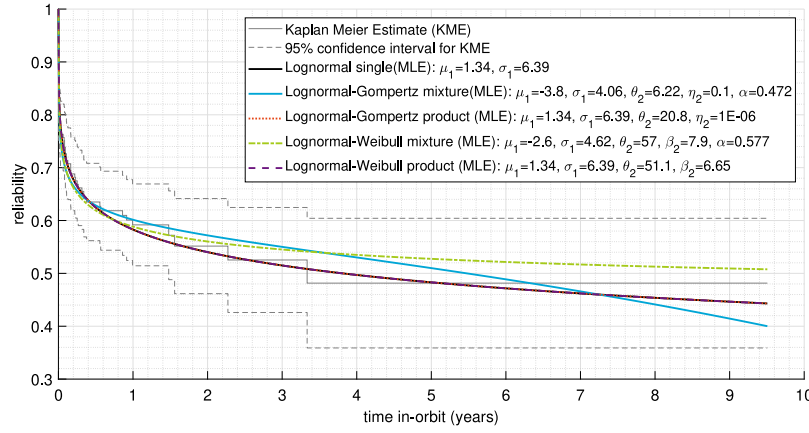


Fig. 5. MLE estimates of the CubeSat failure database.

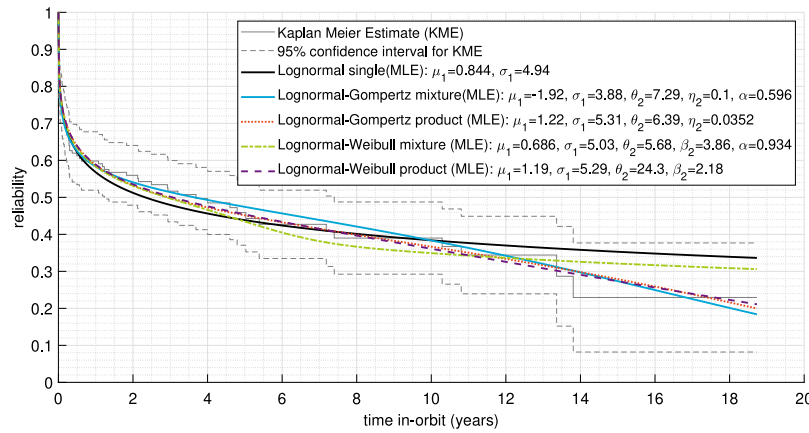


Fig. 6. MLE estimates of the small satellite failure database.

5. Subsystem reliability estimates

Having determined satellite reliability, the next step is to determine subsystem reliability.

5.1. From satellite to subsystem model

Eq. (18) provides the general relation between the system reliability R_{sys} and its subsystems reliability $R_{ss,i}$ for all n subsystems:

$$R_{sys}(t) = \prod_{i=1}^n R_{ss,i}(t). \quad (18)$$

A first subsystem reliability estimate can be made by assuming all $R_{ss,i}$ are identical. This estimate can be used as prior for Bayesian inferences on specific subsystem failure data. The CubeSat failure database contains information on the determined or suspected subsystem that led to a satellite failure: Attitude Determination and Control (ADCS), Command and Data Handling Subsystems (CDHS), Communication Subsystem (COMMS), Structure & Deployment Mechanisms (STS & DepS), Electrical Power Subsystem (EPS) and Payload (P/L). The Thermal Control Subsystem (TCS) is missing in the database because this subsystem is typically passive in CubeSats or embedded in other subsystems. For example, a battery heater would be considered part of the EPS. The electrical interfaces between subsystems are also allocated to the major subsystems (e.g. data interfaces to CDHS, power distribution to EPS). For a study on the effect of redundancy of subsystems, the allocation of failures to these subsystems is not ideal as redundancy is typically applied to physical units and their interfaces. It is expected that more advanced CubeSats may have more physical units and/or

more sophisticated units. For example, an advanced ADCS may comprise an additional board with reaction wheels. If n would represent physical units, its value would differ per satellite. On the other hand, there is a correlation between the sophistication of CubeSats and the experience of its developers as functionality is often added in follow-up satellite projects. Because of the lack of insight of failures related to all these aspects, the potential analysis is limited to the breakdown in aforementioned subsystems. For the research goal of this paper, investigating the impact of subsystem redundancy for CubeSats in general, the limited breakdown is considered to be acceptable. When assessing the reliability of a specific CubeSat, the estimates in this study should be complemented with insights in the complexity of the design, team experience and intensity and results in testing. An example of CubeSat specific reliability estimation and growth is provided by Langer [45].

Besides the breakdown in six subsystems, the database also contains a category ‘unknown’ for satellite failures in which the fatal subsystem is unknown. With 23 out of 71 failures classified as unknown, this is the largest group. Censoring the items would lead to a considerable overestimation of the reliability, so removal of these entries from the database for subsystem analysis is therefore considered to be the best solution. A final check is required if the product of subsystem reliability estimates approximates the general CubeSat reliability estimate.

Eq. (18) can only be used for non-redundant subsystems or the aggregated reliability of redundant subsystems. The database does not contain any information on subsystem redundancy. According to the CubeSat survey on data busses, approximately 10% of the implemented data busses are redundant [11]. While this does not allow to draw any conclusion about redundancy in other subsystems, it may be used as an indication that full redundancy is not widely implemented yet. For none

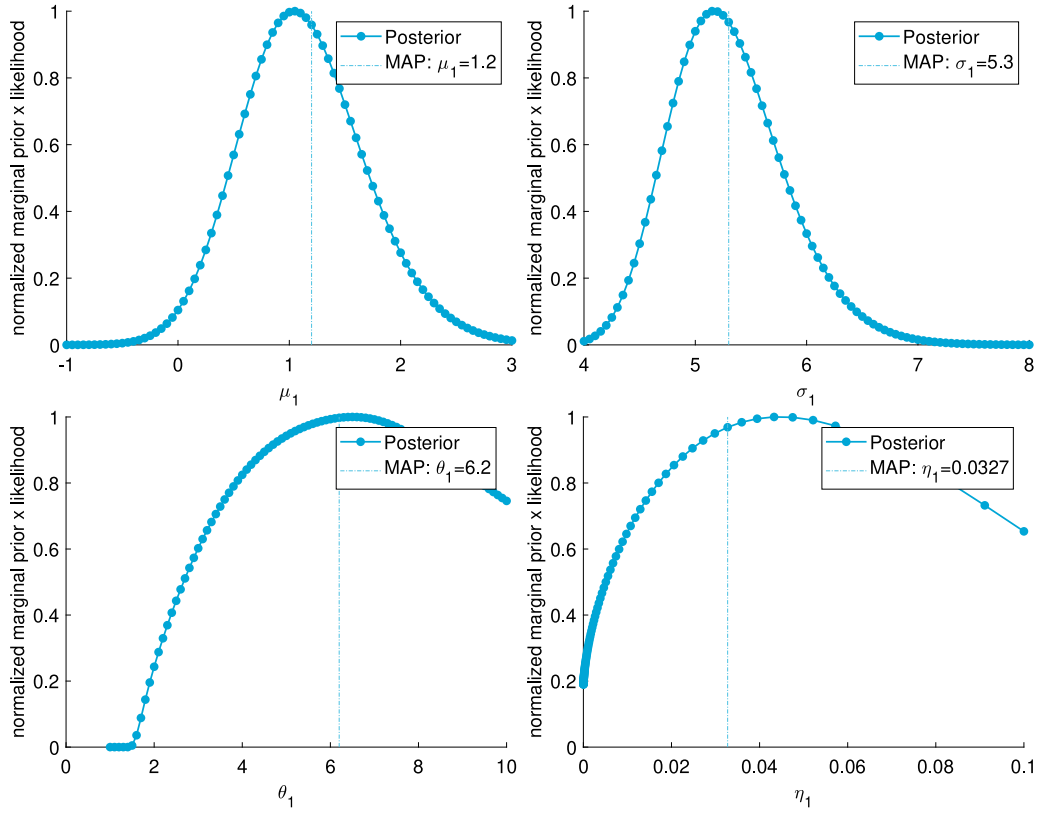


Fig. 7. Marginalized posterior distribution of Lognormal-Gompertz product for small satellite failure database.

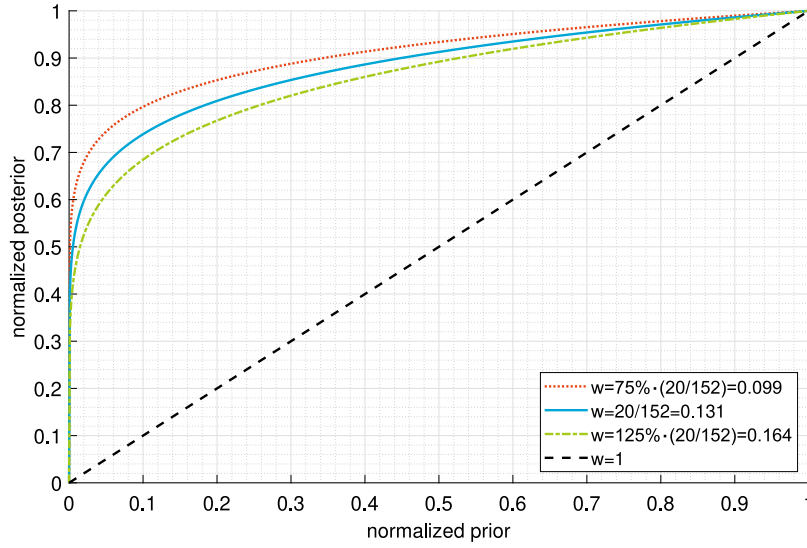


Fig. 8. Posterior-to-prior translation, normalized to MAP.

of the reported satellite failures, a dual failure of a redundant subsystem was mentioned as cause. Furthermore, the impact of redundancy of subsystems on satellite lifetime extension is expected to be highest for the population subject to wear-out, for which the database provides little information as discussed in Section 2.3. For immaturity failures, it is assumed that the vast majority of CubeSat failures are due to single subsystem failures or common mode failures. Eq. (18) is therefore used to estimate the reliability model parameters of non-redundant subsystems.

The posterior for satellites is converted to a prior for subsystems and Bayesian inference is subsequently applied using the specific subsystem

failure data. The posterior for the satellite reliability parameters is translated into an equivalent posterior for subsystems based on identical distributions. In this case Eq. (18) translates into Eq. (19) and, using Eq. (4), into Eq. (20).

$$R_{sat}(t) = R_{ss}(t)^n \quad (19)$$

$$f_{sat}(t) = n \cdot f_{ss}(t) \cdot R_{ss}(t)^{n-1} \quad (20)$$

For immaturity failure *imm.* and the wear-out *w.o.*, subsystem reliability can be split by Eq. (21). This means that the posterior from the results in Section 4.2 can be used as-is by calculating the subsystem

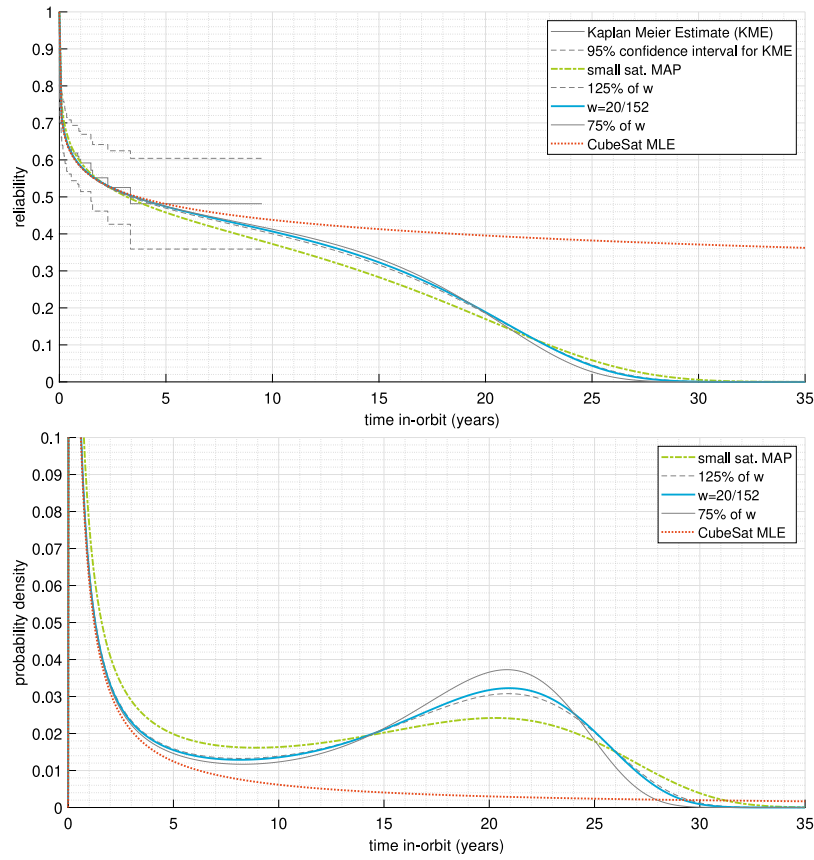


Fig. 9. CubeSat MAP distributions for various priors.

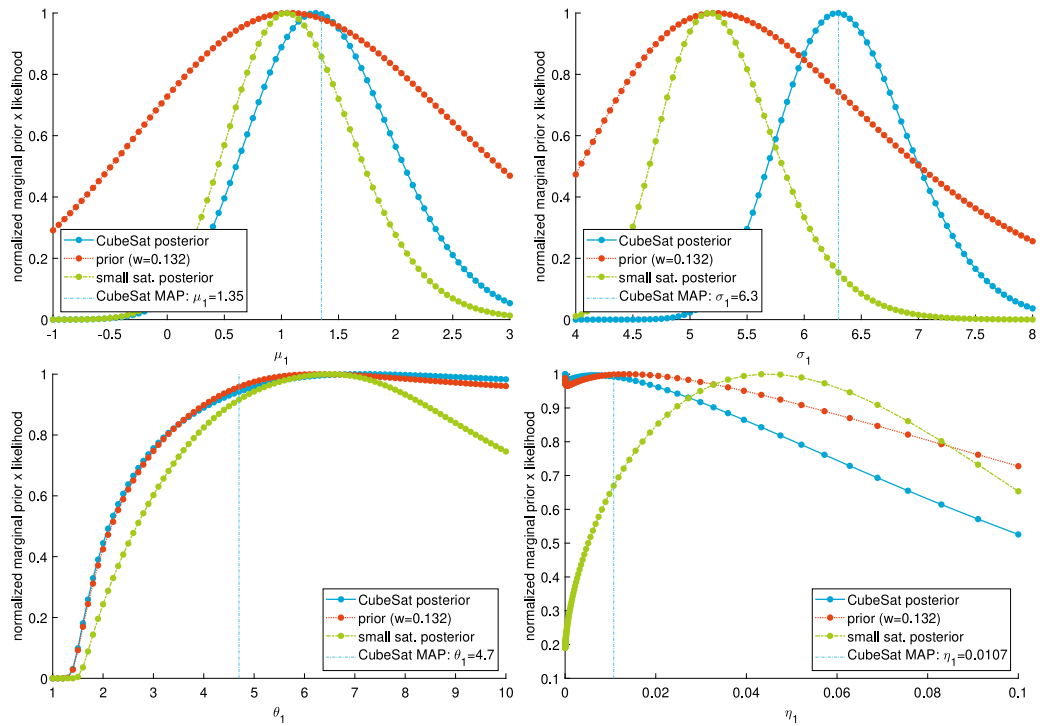


Fig. 10. Marginalized prior and posterior distributions, normalized to MAP.

parameters associated with the satellite parameters for immaturity and wear-out separately.

$$R_{sat}(t) = R_{ss,imm}(t)^n \cdot R_{ss,w.o.}(t)^n \quad (21)$$

For wear-out, Eq. (22) holds when $\theta_{sat} = \theta_{ss}$ and $\eta_{sat} = n \cdot \eta_{ss}$.

$$\exp \left[-\eta_{sat} \left(e^{\left(\frac{t}{\theta_{sat}} \right)} - 1 \right) \right] = \exp \left[-\eta_{ss} \left(e^{\left(\frac{t}{\theta_{ss}} \right)} - 1 \right) \right]^n \quad (22)$$

For immaturity failure, the integral of the Lognormal reliability function cannot be solved in closed form. Instead, the subsystem parameters can be calculated numerically by a discrete representation of the curve for R_{sat} for each set of $(\mu_{sat}, \sigma_{sat})$ in the parameter grid of the satellite posterior. Subsequently, the least squares estimator is used to find the parameters (μ_{ss}, σ_{ss}) for n subsystems for each grid location. This method has been performed using 1000 data points on a logarithmic scale between 0.001 and 100 years. The resulting grid values for μ_{ss} and σ_{ss} values are dependent on both μ_{sat} and σ_{sat} , so subsequent Bayesian inference with subsystem data should be based on the original satellite parameter grid which is then converted to subsystem estimates point-by-point. With this approach, the posterior can be calculated and subsystem MAPs can be found for each point, but the new posterior distribution cannot be marginalized for μ_{ss} and σ_{ss} . The R_{adj}^2 values for each new distribution based on the reliability product of $n = 6$ subsystems with respect to the original reliability of satellites ranges from 0.9991 to 0.9996. This is considered to be a near perfect fit. Fig. 12 shows that the difference between the satellite reliability MAP and the product of the approximated subsystem reliability parameters is indeed small. Using this approach the satellite MAP corresponds to subsystem parameters of $\mu_1 = 13.2$, $\sigma_1 = 9.59$, $\theta_2 = 4.7$ and $\eta_2 = 0.0018$. The converted satellite-to-subsystem posterior will act as prior for Bayesian inference of the specific subsystem data.

5.2. Results of subsystem model estimates

The satellite posterior should be weakened to act as prior, as explained in Section 4.2. In this case a weight of $w = 1/n = 1/6$ is applied for Eq. (17). Using this approach, the failure data will dominate over the prior when there are relative more failures allocated to a subsystem than one-sixth of the satellite failures. If there are less failures for a subsystem, the prior will dominate. Again, the results with $\pm 25\%$ of this weight are also calculated to determine the sensitivity of the results with respect to the weight.

The limited number of failures for each subsystem limits the confidence of the Kaplan–Meier Estimate. Secondly, the large proportion of failures allocated to an unknown subsystem (32%) means that the KME is too optimistic for some subsystems and the lower confidence bound is too high for all. Worst case, if all unknown failures would be allocated to a specific subsystem, the reliability could be 0.32 lower at the end of the observation window. For these reasons, the subsystem KME is not a good reference and the goodness-of-fit loses its meaning and is therefore not provided. The remaining results are presented in Table 4, which provides the Bayesian MAP estimate using $w = 1/6 \pm 25\%$ as well as the MAP of the prior (the parameters for all subsystem distributions identical) and the MLE of the subsystem failure data.

From Table 4 it can be seen that MLE provides unrealistically high values for the wear-out mode $Mode_2$ and the reliability after 50 years $R_{t=50y}$. For the Bayesian estimates using $w = 1/6$ this is all near zero which is more plausible. The results are not significantly sensitive to varying of w while the AIC value is closer to the MLE than to the prior MAP for subsystems with more failure data. Fig. 11 provides the MAP estimates using $w = 1/6$. Using these new estimates, the subsystem product reliability can be calculated and compared to the satellite reliability estimate. This is shown in Fig. 12.

The goodness-of-fit of the product of identical subsystem reliability compared to its original satellite estimate is $R_2 = 0.997$ for the first 10 years and $R_2 = 0.998$ for the first 30 years. The goodness-of-fit of

Table 4

Results of the MAP estimates for Cubesat subsystem failures.

S/S	Estimate	AIC	Mode ₂	R _{t=50y}
ADCS	S/S MLE	8.2	563.4	97.56%
	75% of w	15.7	23.7	0.00%
	w = 1/6	15.9	24.2	0.00%
	125% of w	15.9	24.8	0.00%
	Prior MAP	20.6	29.7	0.00%
CDHS	S/S MLE	72.3	26.7	0.05%
	75% of w	73.8	31.5	0.00%
	w = 1/6	74.0	33.2	0.03%
	125% of w	74.2	34.4	0.14%
	Prior MAP	76.0	29.7	0.00%
COMMS	S/S MLE	28.9	187.5	85.40%
	75% of w	29.0	24.2	0.00%
	w = 1/6	29.0	24.4	0.00%
	125% of w	29.0	24.6	0.00%
	Prior MAP	29.2	29.7	0.00%
STS & DepS	S/S MLE	27.3	2766.6	91.58%
	75% of w	34.3	23.9	0.00%
	w = 1/6	34.8	24.6	0.00%
	125% of w	35.1	25.1	0.00%
	Prior MAP	40.7	29.7	0.00%
EPS	S/S MLE	63.0	187.0	63.77%
	75% of w	64.1	25.6	0.00%
	w = 1/6	64.6	26.4	0.00%
	125% of w	65.0	27.1	0.00%
	Prior MAP	69.5	29.7	0.00%
P/L	S/S MLE	27.3	2766.6	91.58%
	75% of w	34.3	23.9	0.00%
	w = 1/6	34.8	24.6	0.00%
	125% of w	35.1	25.1	0.00%
	Prior MAP	40.7	29.7	0.00%

the product of the individual subsystem reliability estimates compared to the satellite estimate is $R_2 = 0.958$ for 10 years and $R_2 = 0.984$ for 30 years. While the difference in the curves can clearly be seen in Fig. 12, this difference is considered acceptable given the limitations in the subsystem failure data. The LSE of the individual subsystem product is $\mu_1 = 1.51$, $\sigma_1 = 5.50$, $\theta_2 = 3.03$ and $\eta_2 = 0.0014$ which is a perfect fit ($R^2 = 1.00$).

5.3. Generating sample data from subsystem model

For the subsequent steps in modelling, the subsystem MAP estimates from Table 4 with $w = 1/6$ can be used as best estimate. Samples from the posteriors can be used to also include the parameter uncertainty in the simulation. These samples can be obtained by applying the following procedure for each subsystem:

1. Draw a random parameter set $(\mu, \sigma, \beta, \eta)$ from the parameter grid.
2. Calculate the likelihood ratio λ for the sample with respect the MAP.
3. Draw a random number p from a uniform distribution $[0,1]$.
4. Select the parameter set if $\lambda > p$.
5. Repeat steps 1–4 until a defined number of parameter sets have been selected.

6. Modelling subsystem redundancy and failure mitigation

The reliability models for subsystems have been established in Section 5. To be able to create a reliability simulation for the satellite with redundant subsystem and/or improved testing, additional models are needed. The failure dependency between units of a redundant subsystem is investigated first. Subsequently, a model is established for the mitigation of immaturity failures through improved testing.

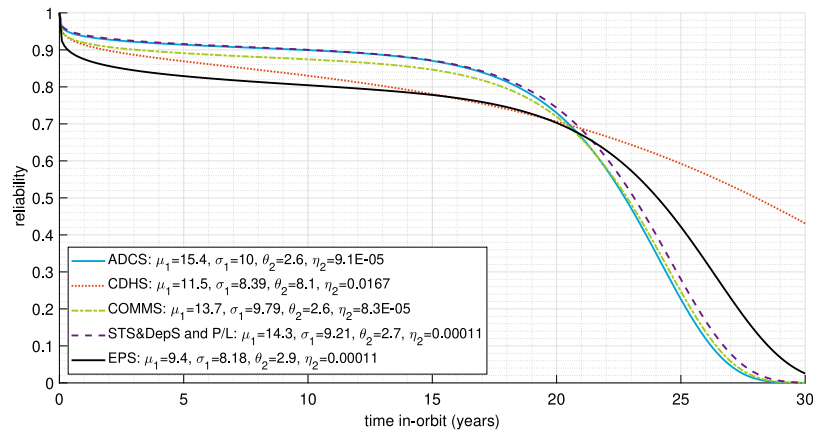
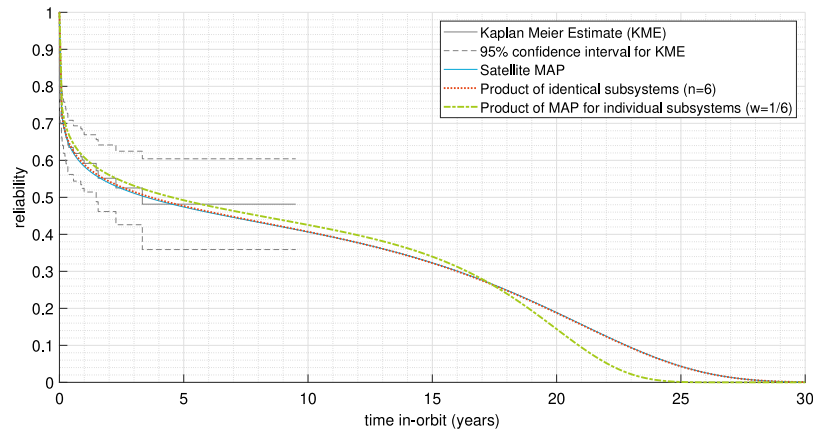
Fig. 11. Reliability MAP estimates using $w = 1/6$.

Fig. 12. Comparison of calculated CubeSat reliability.

Table 5

Definitions of dependent failures as adopted from Borsok [46].

Dependent failure	The likelihood of a set of events, the probability of which cannot be expressed as simple product of unconditional failure probabilities of the individual events.
Common cause failure	This is a specific type of dependent failure that arises in redundant components where simultaneous (or near simultaneous) multiple failures result in the same way or in different ways from a single shared cause.
Common mode failure	This term is reserved for common-cause failures in which multiple items fail in the same way.
Cascade failure	These are all those dependent failures that have no common cause, i.e. they do not affect redundant components.

6.1. Unit dependence of redundant subsystems

Definitions for dependent failures can be confusing. In Table 5 an overview of definitions is provided, where the vertical bars on the left indicates that the underlying failure type(s) are a subset of the above.

Cascade failures are ignored in this study as only fatal failures are considered and all subsystems are assumed to be critical. For the other dependent failures, several methods are available to model the failure dependency between redundant units. Examples are the basic parameter model, the alpha factor model and beta factor model [47]. However, they can be applied only if the dependency is time-independent. An approach is required to address dependent failures for redundant systems over time. Two dependencies of the secondary unit on the primary unit

are considered: the start of the lifetime at risk and the time-to-failure dependency.

The start of the lifetime at risk of the secondary unit depends on whether the underlying failure causes apply only when the unit is operational or also when the unit is switched off. It also depends on whether the secondary unit is hot- or cold-redundant. For CubeSats however, power consumption is a major issue and therefore it is assumed that only cold redundancy can be considered at subsystem level. Failure causes can be differentiated between environmental effects and operational effects, relating to an external and internal cause respectively. For a cold redundant subsystem, which is assumed for this study, the cumulative environmental effects (such as thermal cycling, ionization, externally induced vibration, UV, etc.) affect both units from the time the satellite is deployed into orbit, regardless of its operational state. The redundant unit may in principle already have failed before it is commanded to be turned on. Most single environmental effects (such as latch-up or bit upsets) only applies when the unit is active. The same holds for all operational effects. The parameter ϵ is introduced as the probability of a failure dependent on orbital lifetime, independent of its operational state. This yields that $1 - \epsilon$ is the probability of failure dependent on operational lifetime. The lifetime of the secondary unit in this case starts after the primary has failed based on the assumption of cold redundancy.

The second dependency relates to the time-to-failure for a redundant subsystem. If the redundant subsystem comprises identical units, a subset of common cause and common mode failures related to design flaws (causing immaturity failure) yields a time-to-failure dependency of secondary unit to the primary. If, for example, an electrical power subsystem fails after one orbit because the battery could not handle

the peak temperature as it was never designed for that temperature or the thermal analysis was flawed, it becomes very likely that an identical redundant unit also fails in approximately the same time span. The same is true if it survives for years. Some immaturity failure root causes are however independent between the units of a redundant subsystem. Failures in component production or assembly of random nature are considered in the scope of immaturity failures but do not yield dependencies between identically designed units. The beta factor model [47] can be adopted to account for the ratio of lifetime dependent failures β for immaturity failures, with $\beta = 0$ for fully lifetime independent and $\beta = 1$ for fully dependent failures. While the initial $f(t)$ used for the primary subsystem has a decreasing probability density over time, the updated $f(t)$ is expected to have a narrow log-normal distribution with its mode (peak density) around the failure time of the primary subsystem. The exact parameters of this distribution are however unknown. Assuming a narrow distribution which is approximately symmetric around the mode, the failure time of the secondary unit can be approximated by that of the primary.

When the primary unit fails due to wear-out, this is due to accumulative effects for which failures typically have a high variance. While the distribution of a specific type of wear-out can differ from the overall wear-out distribution, the time-to-failure dependency for a secondary unit is limited and unknown. Therefore, the original probability density used for the primary unit can best be applied to the secondary as well and a potential time-to-failure dependency is ignored.

6.2. Modelling of failure dependencies

The database from the CubeSat survey [11], which is used as input for the CubeSat failure database [10], comprises confirmed or expected root causes of satellite and/or mission failure for 30 of 60 launched CubeSats. Using this input, the cause has been classified in terms of time-to-failure-dependence (β) and its lifetime-at-risk dependence (ϵ) in case a hypothetical identical redundant unit would have been applied. Regarding time-to-failure, this yields 12 dependent, 6 independent and 12 unknown failures. Regarding lifetime-at-risk dependence this yields 13 failures related to orbital lifetime, 10 related to operational lifetime and 7 unknown failures. The problem can be approached as two Bernoulli experiments for (β) and (ϵ), where a 'success' relates to cases which confirm time-to-failure and orbital lifetime dependence respectively and a 'failure' relates to cases which are time-to-failure independent and dependent on operational lifetime, respectively. When applying Bayesian inference on each parameter, the beta distribution can be used as conjugate prior with hyper-parameters a and b [48]. This means that the prior and posterior are both a beta probability density distribution f_{beta} , with a for the number of successes and b for the number of failures.

$$f_{\text{beta}} = \frac{1}{\int_0^1 x^{a-1}(1-x)^{b-1} dx} \cdot x^{a-1}(1-x)^{b-1} \quad (23)$$

Using $a_{\text{prior}} = 1$ and $b_{\text{prior}} = 1$ yields a uniform prior over the range $[0,1]$. The number of 'successes' and 'failures' can be added respectively to obtain the posterior. The classification according to the survey, however, yields numbers of unknown dependencies. Ignoring them would yield a too strong posterior. Therefore the ratio of classified to total failures n_{tot} is applied as weight factor on the classified failures as provided in Eqs. (24) and (25). This yields $a_{\beta} = 8.2$, $b_{\beta} = 4.6$, $a_{\epsilon} = 11.0$ and $b_{\epsilon} = 8.7$ for which the results are shown in Fig. 13. For the simulation, samples from these posteriors will be drawn.

$$a_{\text{post}} = a_{\text{prior}} + \frac{n_s + n_f}{n_{\text{tot}}} n_s \quad (24)$$

$$b_{\text{post}} = b_{\text{prior}} + \frac{n_s + n_f}{n_{\text{tot}}} n_f \quad (25)$$

6.3. Modelling of immaturity failure mitigation

The reduction of the immaturity failures for a satellite without subsystem redundancy is investigated for the case that project resources, otherwise required for implementing subsystem redundancy, are allocated to measures which reduce immaturity failures instead, e.g. increased testing. In a statistical study of CubeSats, Swartwout states that early failures of CubeSats developed at universities can mainly be attributed to insufficient or even complete lack of functional system level testing [4]. From the CubeSat survey it follows that the average duration of testing at fully integrated system level of CubeSats is approximately two months [45].

When applying an extensive test campaign to the system for a period of six months, including fixes and improvements where necessary, it is expected that immaturity failures can be reduced significantly. These test should include duration testing of several months, all possible environmental tests and state-based testing following Failure Mode Effect and Criticality Analysis. The reliability can be further improved if a satellite platform is launched, tested in-orbit and subsequently iterated based on the operational lessons learnt. To model these cases of 'improved testing' and 'iterative development', the failures described in Section 6.2 are analysed for this purpose. All of these failures can be classified as immaturity failures. For each failure the likelihood that improved testing and iterative development respectively would mitigate the failure is approximated. This is done in coarse steps of 0.25, with 1 for almost certainly mitigated and 0 for almost certainly not mitigated. The sum of the likelihoods for the 30 analysed satellites yields an expected mitigation of 16 and 26.5 satellite failures for the improved testing and launched iteration cases respectively. Beta distributions for the likelihood parameter of subsystem immaturity failure mitigation P_{imp} for the improved testing case and P_{iter} for the iterative development case can be constructed in similar fashion to the dependency parameters as explained in Section 6.2. Using a uniform prior, the distribution inputs become $a_{\text{imp}} = 17$, $b_{\text{imp}} = 15$, $a_{\text{iter}} = 27.5$ and $b_{\text{iter}} = 4.5$. The results are shown in Fig. 14.

7. Satellite reliability simulation model

A simulation model is required which uses the estimated subsystem failure probability distributions as input and simulate the results of the satellite reliability with redundant subsystems using the dependency probability distributions and the satellite reliability with improved testing using the mitigation probability distributions. To set up this simulation, several tools have been considered: fault trees, event trees, Petri-nets, Markov Chains and activity flows. The fault tree is a limited representation for a time-dependent model as it does not include the distribution over time and the impact of the dependencies on this distribution. In fault trees, reliability values are typically limited to mutually independent failure probabilities at a given time. The use of dynamic fault trees [49] or Petri-nets [50] could be considered for this purpose as these modelling tools provide options to introduce the impact of these dependencies. However, these tools are focused on multi-level failures which is not in the scope of this study. For the intended model, their diagrams are not as easy to interpret, so a more simple representation is desired. Event trees, although typically used to analyse cascading effects of initial (failure) events, provides a decision-tree structure [51] which can introduce the redundant subsystem failure dependencies and failure mitigation. This would however lead to a high number of branches for the intended simulation and would still require adaptation to introduce the time-dependent models. Markov Chains have been used for CubeSat reliability modelling by Engelen et al. [52]. For simulations of limited complexity they can be used by replacing the typical fixed value probabilities with time dependent models. Like event trees, they will however become very large. Moreover, these tools cannot be used to show all steps required to perform a Monte Carlo simulation to be able to run the different

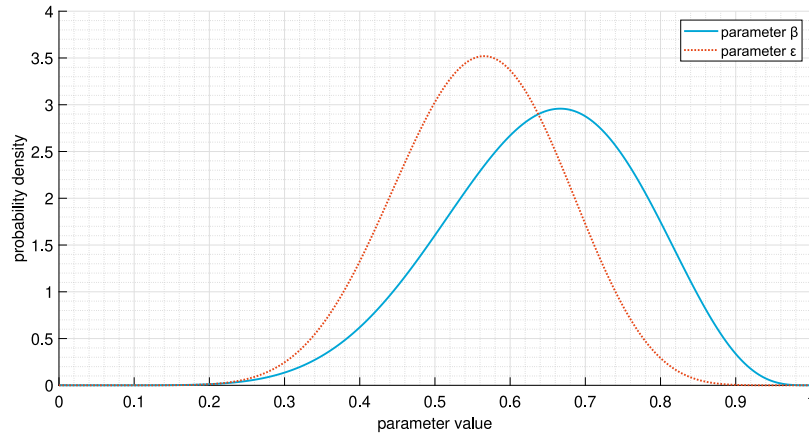


Fig. 13. Posterior distribution of dependence parameters β and ϵ .

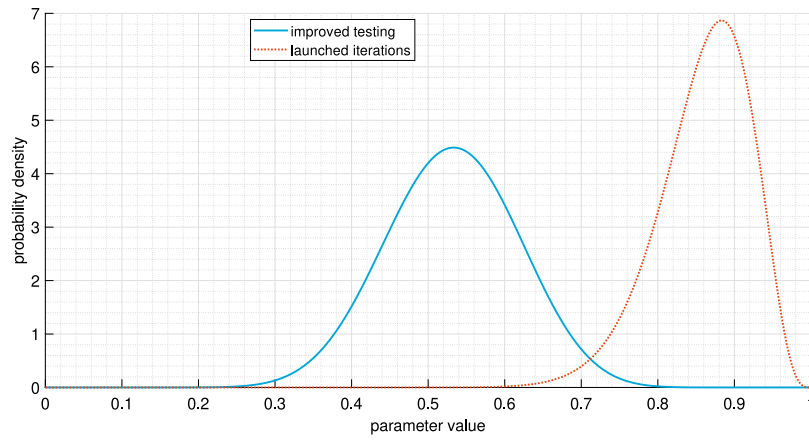


Fig. 14. Posterior distribution of mitigation parameter P_{imp} and P_{iter} .

cases. A ‘reliability modelling flow’ is therefore used as a new type of representation for dependent binary-state failures. It is based on an Universal Modelling Language (UML) activity flow. The full simulation model is provided in Fig. 15, where the two high level blocks in bold are worked out at lower level in Figs. 16 and 17.

Fig. 15 provides the modelling flow for the reliability of satellites. For all $n_{S/S}$ subsystems, samples are generated. The lifetime of the satellite with redundancy $t_{red.(sat)}$ and without redundancy $t_{single(sat)}$ is modelled as the minimum lifetime of all of its subsystems. A number of $n_{samples}$ satellite samples are generated for each simulation. When enough of such output samples are randomly generated, the distribution of the output is representative for satellites with and without redundant subsystems. Given a sufficiently high number of samples, a Kaplan–Meier Estimator (KME) can be used. Parametric estimates of the output are not needed as they will result to the same figure. A number of n_{sim} simulation runs are performed using the parameter samples from the posterior (see Section 5) and samples of the dependency parameters (see Section 6.2) to perform a full Monte Carlo simulation which includes the uncertainties on the input parameters. Each simulation results in a different estimated reliability curve.

The model for a subsystem failure is provided in Fig. 16. This flow uses sampled input parameters for the subsystem Lognormal-Gompertz product PDF (μ , σ , θ , η) and the dependencies (β , ϵ). It first creates intermediary failure time values t_1 to t_4 for the PDF and reference probabilities p_β and p_ϵ for the dependencies. The failure time values t_1 to t_4 can be generated by using a random generator for a uniform distribution over $[0, 1]$ to create samples for $F(t)$ which are subsequently put into the inverse transform of $F(t)$. An example for the Gompertz distribution is provided by Eq. (26). To generate a sample

for the Lognormal-Gompertz product, the minimum of the samples of the Lognormal distribution and Gompertz distribution for both units is taken as remaining sample. Using these values and following the dependence decisions in the flow, a sample for the primary unit $t_{prim.}$ and the secondary unit $t_{sec.}$ is calculated and the subsystem failure time $t_{red.}$ for a redundant subsystem is generated as output. The primary unit sample $t_{prim.}$ is the output for a subsystem without redundancy.

$$t = \theta \cdot \ln \left[1 - \frac{\ln(1 - F(t))}{\eta} \right] \quad (26)$$

When the simulation provides estimates based on all samples of the original subsystem failure distributions (Section 5.2), this is called the ‘reference case’. For the cases of ‘improved testing’ and ‘iterative development’, some of the subsystem failures due to immaturity are mitigated. The activity flow for failure mitigation is provided in Fig. 17. Probability values p_1 and p_2 are generated and compared against the samples mitigation parameters P_{imp} and P_{iter} . The mitigation parameters change for each simulation of $n_{samples}$ and are taken from their posterior beta distributions explained in Section 6.3. In case of redundant subsystems, mitigation can only apply to secondaries which fail independently due to immaturity. If the root-cause subsystem failure is mitigated, this failure time is flagged as censor time instead which can be used as input for the KME.

8. Results of the satellite reliability simulation

The reliability of a satellite with redundant subsystems is simulated in 100 runs, each producing 10,000 samples for satellite failure using the model explained in Section 7. The results for the reliability over

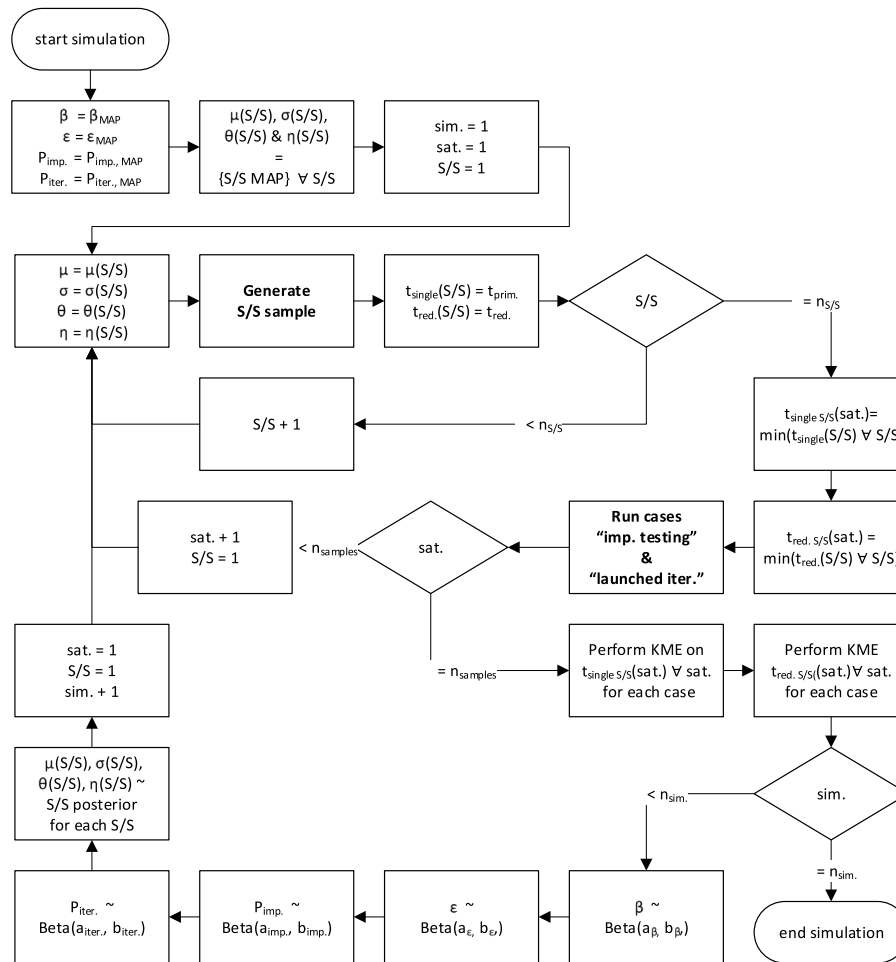


Fig. 15. Reliability modelling flow for satellite simulation.

Table 6
Satellite simulation reliability at specified time in-orbit.

Subsystems	Case	Reliability mean of simulations			
		1 y	3 y	5 y	10 y
Non-redundant	Reference	0.60	0.52	0.49	0.42
	Reference	0.73	0.67	0.64	0.59
Non-redundant	imp. testing	0.79	0.74	0.71	0.65
	imp. testing	0.85	0.81	0.79	0.76
Non-redundant	Iterative dev.	0.93	0.91	0.89	0.85
	Iterative dev.	0.93	0.91	0.90	0.89

time is provided in Fig. 18 in the form of a Kaplan–Meier Estimate. The thick lines represent the reliability over time using the subsystem MAP estimates as input, while the smaller lines are the results from simulation runs using other samples of the subsystem posterior distributions as explained in Section 5. In Table 6 the reliability after a specified time in orbit is provided for the MAP inputs and the mean of all simulation runs.

Fig. 18 already clearly shows a ranked order based on the subsystem MAP inputs for the first 15 years in orbit. A similar ranking can be seen in Table 6 which provides the mean reliability of all simulation runs after 1, 3, 5 and 10 years. A reserved conclusion would be that improved testing and iterative development yields a major improvement for these failures. The lines for the different simulations in Fig. 18, however, partially overlap and could therefore in specific simulation runs yield a different outcome. For this purpose, the reliability for each simulation run after 1, 3, 5 and 10 years are compared. Most CubeSat design lifetimes will be in the range of one and five years (see Section 2.3). The

reliability at 10 years is therefore of limited interest, but added for some unique future missions. In Fig. 19 a scatter plot is provided, with on the horizontal axis the reliability for a satellite with redundant subsystems and on the vertical axis the reliability for a satellite without subsystems but improved testing instead. The diagonal line indicates the theoretical boundary where the reliability of the two options would be equal. This figure confirms that allocating resources to improved testing in general has a better impact on reliability than subsystem redundancy. Only for very long missions of 10 years, there is a small chance that subsystem redundancy is superior to improved testing.

It is also interesting to see the effect of combining subsystem redundancy with improved testing, compared to a non-redundant satellite with the same level of testing. Fig. 20 shows a significant improvement can be expected with redundant subsystems. This does however require substantial project resources, especially for unique single satellite missions.

Finally, it is interesting to compare satellites with and without redundant subsystems in the case of iterative development. This is especially interesting for CubeSat networks or a standardized CubeSat platform for multiple missions. In this case there is not so much of a trade to make based on project resources, as iterations may already been foreseen by programmatic choice. The results are provided in Fig. 21. For missions up to 3 years, the results are scattered around the equality line. Only for missions of 5 years or more, redundancy does pay off.

9. Conclusions

The answer to the question “What leads to higher CubeSat reliability over its mission life time: full subsystem redundancy or improved

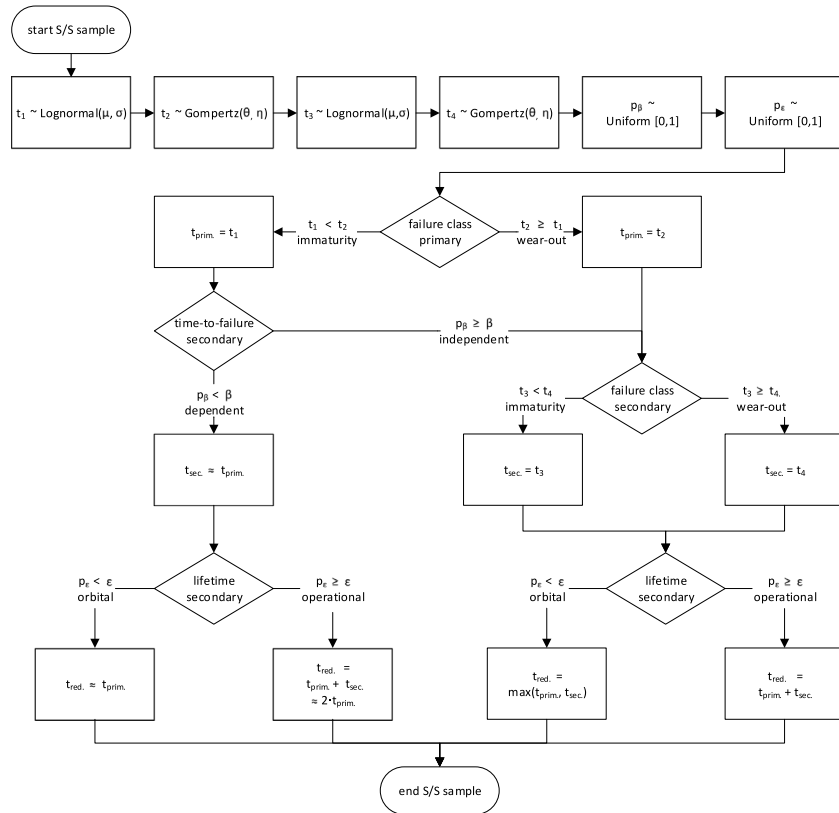


Fig. 16. Reliability modelling flow for a (cold redundant) subsystem.

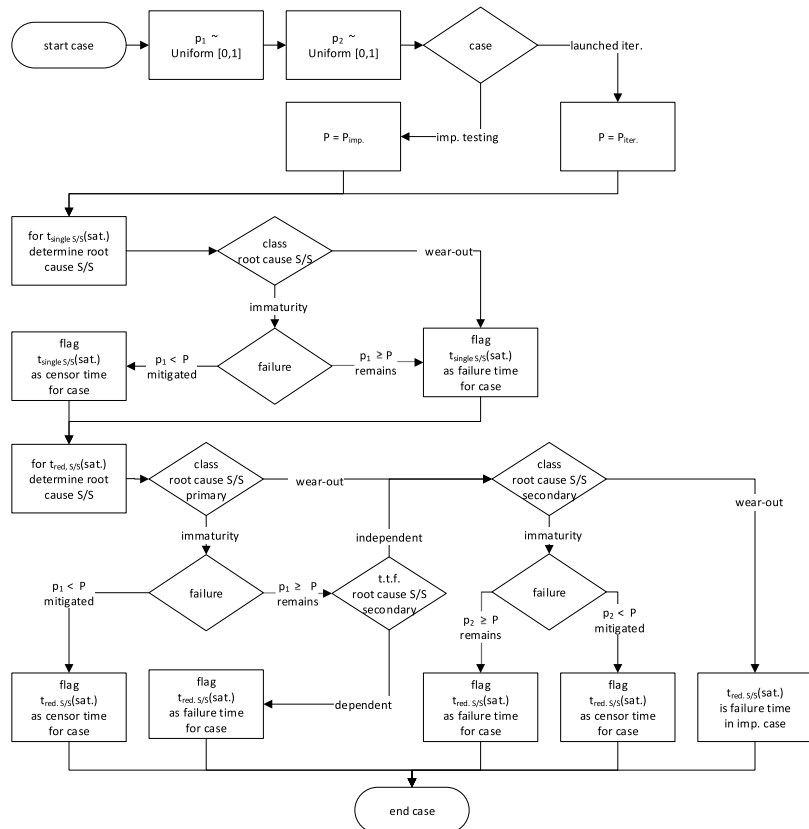


Fig. 17. Reliability modelling flow for failure mitigation.

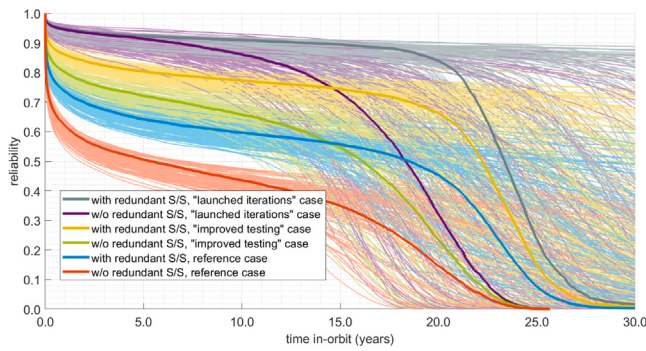


Fig. 18. Kaplan-Meier Estimate simulation results.

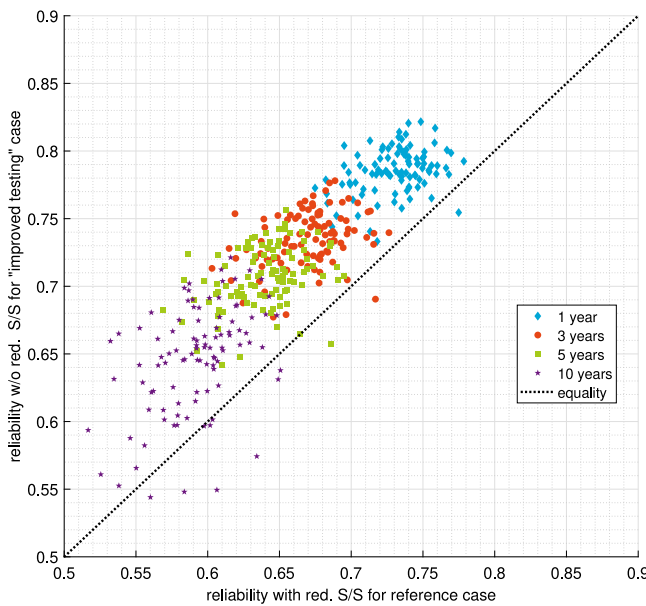


Fig. 19. Satellite reliability scatter for redundant subsystems vs. non-redundant subsystems with improved testing.

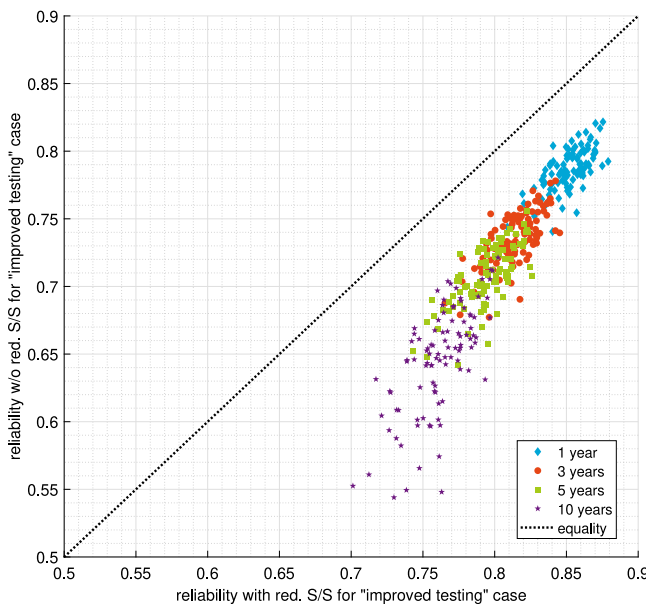


Fig. 20. Satellite reliability scatter for redundant vs. non-redundant subsystems, both with improved testing.

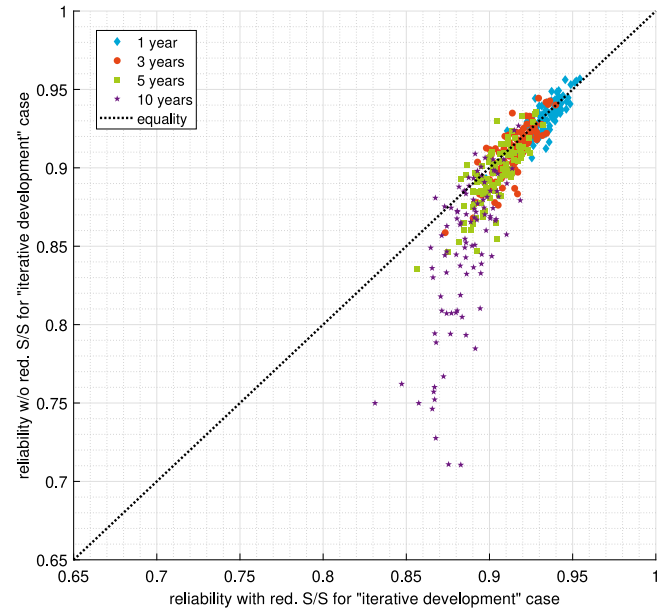


Fig. 21. Satellite reliability scatter for redundant vs. non-redundant subsystems, both with iterative development.

testing?” is that improved testing yields the best results for most missions, based on the simulation results presented in Fig. 19. It has the additional benefit to increase potential payload volume and has a lower platform cost compare to a satellite with redundant subsystems. Furthermore, iterative development is the best strategy for series and networks of CubeSats and/or bus platforms. A satellite with redundant subsystems remains more reliable than a satellite without redundancy above 10–15 years as shown in Fig. 18. This is however beyond the typical useful lifetime of CubeSats and therefore considered to be of limited importance. It can also be seen, by the spread of the simulated curves, the uncertainty of the model increases over time which is mainly due to limited observations causing relatively large uncertainties of the wear-out parameters.

A Lognormal-Gompertz product provides the best parametric reliability model for a CubeSat and small satellites in general. The Kaplan-Meier estimator is a necessary step to show the pure satellite observation data and provide a reference for parametric models. Maximum likelihood estimators are good for model comparison, but lack the ability to introduce prior knowledge. Furthermore, they do not provide a full posterior which limits the ability to properly model uncertainties in subsequent modelling. Bayesian inference is the best approach to overcome these limitations. This paper provides an example of the use of these tools, which can be of interest to satellite reliability modelling or even reliability modelling in general. The satellite simulation model as introduced in Section 7 is an innovative method which can be applied to other satellite size categories as well as other complex systems.

The research question implies a choice between allocating additional resources to either implement redundant subsystems or to improve testing on a satellite without redundancy. Such a choice may not be applicable if resources allow for both improvements. In this case, the reliability of a satellite with redundancy has a significantly higher reliability over time compared to a satellite without. However, applying redundancy does not only consume additional organizational resources. It also consumes a considerable amount of volume of the satellite which leaves less room for the payload. Moreover, this study assumes a flawless failure detection, isolation and recovery mechanism which arbitrates between the redundant units of a subsystem which may be too optimistic in reality. For a single satellite in a project with limited

resources, it is therefore considered to be a better strategy to aim for reduction of immaturity failures through extensive testing. For satellite networks, ‘swarm robustness’ could be achieved [52] and individual satellite losses may be acceptable. Only for single satellite missions longer than 10 years or operating in harsher environments beyond LEO, redundant subsystems may be required to improve reliability to acceptable values.

Immaturity failures can be reduced by iterating on the satellite bus platform with subsequent launches and high reliability can be achieved, as shown in Fig. 21. This would extend improved on-ground testing to in-orbit testing. A condition for this approach is that improvements in performance of each subsequent design remain limited and improvements are primarily focused on the reliability of the design. A modular philosophy of CubeSats where subsystems from different manufacturers are procured and integrated is not compatible with this strategy as the lack of direct involvement may prohibit the required improvement or the selection of a different model can introduce new immaturity failure risks at subsystem or satellite level. The entire iterative platform development must therefore be under control of one party or consortium. While this may have its limitations, it opens the possibility to introduce advanced architectural concepts which deviate from the modular approach [25]. An example of such architecture is the integration of satellite core functionality into a single physical unit to reduce its effective volume and reduce the component count (potentially increasing reliability further). Another example is the use of advanced outer panels which reduce wiring harness and integration complexity.

CRedit authorship contribution statement

J. Bouwmeester: Conceptualization, Methodology, Software, Formal analysis, Writing – original draft. **A. Menicucci:** Supervision, Writing – review & editing. **E.K.A. Gill:** Supervision, Writing – review & editing.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgements

The authors would like to thank Martin Langer, Jian Guo and Liora Monas for sharing their failure databases for this study.

References

- [1] Puig-Suari J, Turner C, Ahlgren W. Development of the standard CubeSat deployer and a CubeSat class picosatellite. In: IEEE aerospace conference proceedings. 2001, <http://dx.doi.org/10.1109/aero.2001.931726>.
- [2] Tafazoli M. A study of on-orbit spacecraft failures. Acta Astronaut 2009;64(2–3):195–205. <http://dx.doi.org/10.1016/j.actaastro.2008.07.019>.
- [3] Guo J, Monas L, Gill E. Statistical analysis and modelling of small satellite reliability. Acta Astronaut 2014;98:97–110. <http://dx.doi.org/10.1016/j.actaastro.2014.01.018>.
- [4] Swartwout M. The first one hundred CubeSats : A statistical look. J Small Satell 2013;2(2):213–33.
- [5] Swartwout M, Jayne C. University-class spacecraft by the numbers: Success, failure, debris. (bot mostly success.). In: Proceedings of 30th aiaa/usu conference on small satellites. AIAA; 2017.
- [6] Kaminskiy M. CubeSat data analysis. Tech. rep., NASA GSFC; 2015.
- [7] Swartwout M. CubeSat Database, 0000. URL <https://sites.google.com/a/slu.edu/swartwout/cubesat-database/census>.
- [8] Villela T, Costa C, Brandao A, Bueno F, Leonardi R. Towards the thousandth CubeSat: A statistical overview. Int J Aerosp Eng 2019.
- [9] Villela, et al. CGEE CubeSat Database. 2018, URL <https://www.cgee.org.br/web/observatorio-espacial/bancos-de-dados>.
- [10] Langer M, Boumeester J. Reliability of CubeSats – statistical data, developers’ beliefs and the way forward. In: Proceedings of 30th aiaa/usu conference on small satellites. Logan: AIAA; 2016.
- [11] Bouwmeester J, Langer M, Gill E. Survey on the implementation and reliability of CubeSat electrical bus interfaces. CEAS Space J 2017;9(2). <http://dx.doi.org/10.1007/s12567-016-0138-0>.
- [12] Kulu E. NanoSats database, 0000. URL www.nanosats.eu.
- [13] Eickhoff J. The flip microsatellite. Springer; 2016, <http://dx.doi.org/10.1007/978-3-319-23503-5>.
- [14] Bayer T. Planning for the un-plannable: Redundancy, fault protection, contingency planning and anomaly response for the mars reconnaissance orbiter mission. In: Proceedings of the space conference & exposition. Long Beach: AIAA; 2007.
- [15] Bouwmeester J, Aalbers G, Ubbels W. Preliminary mission results and project evaluation of the delfi-c3 nano-satellite. In: Proceedings of the 4s symposium. Rhodes: ESA; 2008.
- [16] Bouwmeester J, Rotthier L, Schuurbiers C, Wieling WT, Horn GVD, Stelwagen F, Timmer E, Tijssen M. Preliminary results of the delfi-n3xt mission. In: Proceedings of the 4S Symposium. Mallorca: ESA; 2014.
- [17] Menchinelli A, Ingiosi F, Pamphili L, Marzioli P, Patriarca R, Costantino F, Piergentili F. A reliability engineering approach for managing risks in CubeSats. Aerospace 2018;5(4). <http://dx.doi.org/10.3390/aerospace5040121>.
- [18] Pantoji GS, Bhat MH, Gwalani PN, Bhulokam AM. Development of a risk management plan for RVSAT-1, a student-based CubeSat program. In: IEEE aerospace conference proceedings, Vol. 2021-March. 2021, <http://dx.doi.org/10.1109/AERO50100.2021.9438156>.
- [19] Nieto-Peroy C, Emami MR. CubeSat Mission: From design to operation. Appl Sci (Switzerland) 2019;9(15). <http://dx.doi.org/10.3390/app9153110>.
- [20] Venturini C, Braun B, Hinkley D, Berg G. Improving mission success of CubeSats. In: Proceedings of the 32nd annual aiaa/usu workshop on small satellites. 2018.
- [21] Doyle M, Dunwoody R, Finneran G, Murphy D, Reilly J, Thompson J, Walsh S, Erkal J, Fontanesi G, Mangan J, Marshall F, Salmon L, Ha L, Mills A, Palma D, de Faioite D, Greene D, Martin-Carrillo A, McKeown D, O’Connor W, Stanton K, Uliyanov A, Wall R, Hanlon L. Mission testing for improved reliability of CubeSats. In: Proceedings of the international conference on space optics. 2021, <http://dx.doi.org/10.1117/12.2600305>.
- [22] Berthoud L, Swartwout M, Blvd L, Louis S, Cutler J, Klumpar D. University CubeSat project management for success. In: Proceedings of the 33rd aiaa/usu conference on small satellites. 2019.
- [23] Bouwmeester J, Santos N. Analysis of the distribution of the electrical power in CubeSats. In: Proceedings of the 4s symposium. Valetta: ESA; 2014.
- [24] Bouwmeester J, van der Linden SP, Povalac A, Gill EK. Towards an innovative electrical interface standard for PocketQubes and CubeSats. Adv Space Res 2018;62(12):3423–37. <http://dx.doi.org/10.1016/j.asr.2018.03.040>.
- [25] Bouwmeester J, Gill E, Speretta S, Uludag S. A new approach on the physical architecture of CubeSats & PocketQubes. J Br Interplanet Soc 2018;71(7):1–13.
- [26] ECSS-Q-ST-30C Rev.1: Dependability. Tech. rep. 3rd, rev.1, European Cooperation for Space Standardization; 2017, p. 1–65.
- [27] Klutke GA, Kiessler PC, Wortman MA. A critical look at the bathtub curve. IEEE Trans Reliab 2003;52(1):125–9. <http://dx.doi.org/10.1109/TR.2002.804492>.
- [28] Wong KL. The physical basis for the rollercoaster hazard rate curve for electronics. Qual Reliab Eng Int 1991;7(6):489–95. <http://dx.doi.org/10.1002/qre.4680070609>.
- [29] Bouwmeester J, Langer M. Database: results of cubesat survey on electrical interfaces & reliability. 2016, <http://dx.doi.org/10.4121/uuid:591ff8f8-b495-4d1c-9c5c-69f6f85ace78>.
- [30] Kaplan EL, Meier P. Nonparametric estimation from incomplete observations. J Amer Statist Assoc 1958;53(282):457–81. <http://dx.doi.org/10.1080/01621459.1958.10501452>.
- [31] Castet J-F, Saleh JH. Satellite reliability: Statistical data analysis and modeling. J Spacecr Rockets 2009;46(5):1065–76. <http://dx.doi.org/10.2514/1.42243>.
- [32] Castet JF, Saleh JH. Satellite and satellite subsystems reliability: Statistical data analysis and modeling. Reliab Eng Syst Saf 2009;94(11):1718–28. <http://dx.doi.org/10.1016/j.res.2009.05.004>.
- [33] Greenwood M. A Report on the Natural Duration of Cancer.. Reports on Public Health and Medical Subjects (Ministry of Health) 33, 1926.
- [34] EPSMA. Guidelines to understanding reliability prediction. Report, 2004.
- [35] Castet JF, Saleh JH. Single versus mixture Weibull distributions for nonparametric satellite reliability. Reliab Eng Syst Saf 2010;95(3):295–300. <http://dx.doi.org/10.1016/j.res.2009.10.001>.
- [36] Liu X. Survival analysis: Models and applications. Wiley; 2012, <http://dx.doi.org/10.1002/9781118307656>.
- [37] Lee ET, Go OT. Survival analysis in public health research. Annu Rev Public Health 1997;18:105–34. <http://dx.doi.org/10.1146/annurev.publhealth.18.1.105>.
- [38] Weibull W. A statistical distribution function of wide applicability. J Appl Mech 1951;18(3):293–7.
- [39] Gompertz B. On the nature of the function expressive of the law of human mortality, and on a new mode of determining the value of life contingencies.. Proc R Soc Lond 1833;2. <http://dx.doi.org/10.1098/rspl.1815.0271>.
- [40] Bayes T, Price R. An essay towards solving a problem in the doctrine of chances. Philos Trans (1683-1775) 1763;53:370–418.

- [41] Dubos GF, Castet JF, Saleh JH. Statistical reliability analysis of satellites by mass category: Does spacecraft size matter? *Acta Astronaut* 2010;67(5–6). <http://dx.doi.org/10.1016/j.actaastro.2010.04.017>.
- [42] Steel RGD, Torrie JH. *Principles and procedures of statistics with special reference to the biological sciences*. McGraw-Hill; 1960.
- [43] Miles DM, Mann IR, Ciurzynski M, Barona D, Narod BB, Bennest JR, Pakhotin IP, Kale A, Bruner B, Nokes CD, Cupido C, Haluza-DeLay T, Elliott DG, Milling DK. A miniature, low-power scientific fluxgate magnetometer: A stepping-stone to CubeSatellite constellation missions. *J Geophys Res Space Phys* 2016;121(12):11839–60. <http://dx.doi.org/10.1002/2016JA023147>.
- [44] Aho K, Derryberry D, Peterson T. Model selection for ecologists: The worldviews of AIC and BIC. *Ecology* 2014;95(3):631–6. <http://dx.doi.org/10.1890/13-1452.1>.
- [45] Langer M. Reliability assessment and reliability prediction of cubesats through system level testing and reliability growth modelling (Ph.D. thesis), Technical University of Munich; 2018.
- [46] Börcsök J, Schaefer S, Ugljesa E. Estimation and evaluation of common cause failures. In: Proceedings of the 2nd international conference on systems. 2007, <http://dx.doi.org/10.1109/ICONS.2007.25>.
- [47] Jones HW. Common cause failures and ultra reliability. In: Proceedings of the 42nd international conference on environmental systems. 2012, <http://dx.doi.org/10.2514/6.2012-3602>.
- [48] Fink D. A compendium of conjugate priors. Tech. rep., Bozeman: Montana State University; 1997.
- [49] Volovoi V. Modeling of system reliability Petri nets with aging tokens. *Reliab Eng Syst Saf* 2004;84(2):149–61. <http://dx.doi.org/10.1016/j.res.2003.10.013>.
- [50] Čepin M, Mavko B. A dynamic fault tree. *Reliab Eng Syst Saf* 2002;75(1):83–91. [http://dx.doi.org/10.1016/S0951-8320\(01\)00121-1](http://dx.doi.org/10.1016/S0951-8320(01)00121-1).
- [51] Andrews JD, Dunnett SJ. Event-tree analysis using binary decision diagrams. *IEEE Trans Reliab* 2000;49(2). <http://dx.doi.org/10.1109/24.877343>.
- [52] Engelen S, Gill E, Verhoeven C. On the reliability, availability, and throughput of satellite swarms. *IEEE Trans Aerosp Electron Syst* 2014;50(2):1027–37. <http://dx.doi.org/10.1109/TAES.2014.120711>.