

Estimating Probability of Failure of a Complex System Based on Partial Information about Subsystems and Components, with Potential Applications to Aircraft Maintenance

Christelle Jacob¹, Didier Dubois², Janette Cardoso¹, Martine Ceberio³, and Vladik Kreinovich³

¹ Institut Supérieur de l'Aéronautique et de l'Espace (ISAE), DMIA department, Campus Supaéro, 10 avenue Édouard Belin, Toulouse, France
jacob@isae.fr, cardoso@isae.fr

² Institut de Recherche en Informatique de Toulouse (IRIT), ADRIA department, 118 Route de Narbonne 31062 Toulouse Cedex 9, France
dubois@irit.fr

³ University of Texas at El Paso, Computer Science Dept., El Paso, TX 79968, USA
mceberio@utep.edu, vladik@utep.edu

Abstract. In many real-life applications (e.g., in aircraft maintenance), we need to estimate the probability of failure of a complex system (such as an aircraft as a whole or one of its subsystems). Complex systems are usually built with redundancy allowing them to withstand the failure of a small number of components. In this paper, we assume that we know the structure of the system, and, as a result, for each possible set of failed components, we can tell whether this set will lead to a system failure. In some cases, for each component A , we know the probability $P(A)$ of its failure; in other cases, however, we only know the lower and upper bounds $\underline{P}(A)$ and $\overline{P}(A)$ for this probability. Sometimes, we only have expert estimates for these probabilities, estimates that can be described as fuzzy numbers.

Usually, it is assumed that failures of different components are independent events, but sometimes, we know that they are dependent – in which case we usually do not have any specific information about their correlation. Our objective is to use all this information to estimate the probability of failure of the entire the complex system. In this paper, we describe methods for such estimation.

Keywords: complex system, probability of failure, interval uncertainty, fuzzy uncertainty, stochastic dependence

1 Formulation of the Problem

It is necessary to estimate the probability of failure for complex systems. In many practical situations applications, we need to estimate the probability of failure of a complex system. The need for such estimates come from the fact

that in practice, while it is possible (and desirable) to minimize the risk, it is not possible to completely eliminate the risk: no matter how many precautions we take, there are always some very low probability events that can potentially lead to a system's failure. All we can do is to make sure that the resulting probability of failure does not exceed the desired small value p_0 . For example, the probability of a catastrophic event is usually required to be at or below $p_0 = 10^{-9}$.

For example, in aircraft design and maintenance, we need to estimate the probability of a failure of an aircraft as a whole and of its subsystems. On the design stage, the purpose of this estimate is to make sure that this probability of failure does not exceed the allowed probability p_0 . On the maintenance stage, we perform this estimate to decide whether a maintenance is needed: if the probability of failure exceeds p_0 , this means that we need to perform some maintenance to decrease this probability to the desired level p_0 (or below).

Information available for estimating system's probability of failure: general description. Complex systems consist of subsystems, which, in turn, consists of components (or maybe of sub-subsystems which consist of components). So, to estimate the probability of failure of a complex system, we need to take into account:

- when the failure of components and subsystems lead to the failure of the complex system as a whole,
- how reliable are these components and subsystems, and
- are the component failures independent events or they are caused by a common cause, and

When the failure of components and subsystems lead to the failure of the complex system as a whole? Complex systems are usually built with redundancy allowing them to withstand the failure of a small number of components. Usually, we know the structure of the system, and, as a result, for each possible set of failed components, we can tell whether this set will lead to a system failure. So, in this paper, we will assume that this information is available.

How reliable are components and subsystems? What do we know about the reliability of individual components? For each component A , there is a probability $P(A)$ of its failure. When we have a sufficient statistics of failures of this type of components, we can estimate this probability as the relative frequency of cases when the component failed. In some case, we have a large number of such cases, and as a result, the frequency provides a good approximation to the desired probability – so that, in practice, we can safely assume that we know the actual values of these probabilities $P(A)$.

In some case, we only have a few failure cases, not enough to get an accurate estimate for $P(A)$. In this case, the only information that we can extract from the observation is the interval $\mathbf{P}(A) = [\underline{P}(A), \bar{P}(A)]$ that contains the actual (unknown) value of this probability.

This situation is rather typical for aircraft design and maintenance, because aircrafts are usually built of highly reliable components – at least the important

parts of the aircraft are built of such components – and there are thus very few observed cases of failure of these components.

In some cases, especially on the design stage, we do not yet have failure statistics, so we have to rely on expert estimates, estimates based on the expert’s experience with similar components. These estimates are usually formulated by using words from natural language, like “about 1%”. A natural way to describe such expert estimates is to use fuzzy techniques (see, e.g., [9, 14]), i.e., to describe each such estimate as a *fuzzy number* $\mathcal{P}(A)$. A fuzzy number means, crudely speaking, that for different degrees of uncertainty α , we have an interval $\mathcal{P}(A, \alpha)$ that contains the actual (unknown) probability $P(A)$ with this degree of uncertainty: e.g., the interval $\mathcal{P}(A, 0)$ contains $P(A)$ with guarantee (uncertainty 0), while the interval $\mathcal{P}(A, 0.5)$ contains $P(A)$ with uncertainty 0.5.

Are the component failures independent events or they are caused by a common cause? In many practical situations, failures of different components are caused by different factors. For example, for an aircraft, possible failures of mechanical subsystems can be caused by the material fatigue, while possible failures of electronic systems can be caused by the interference of atmospheric electricity (e.g., when flying close to a thunderstorm). As a result, usually, it is assumed that failures of different components are independent events

However, sometimes, we know that the failures of different components are caused by a common cause. In this case, the failures of different components are no longer independent events. In such situations, often, do not have any specific information about the correlation between these failures. This is a typical situation in aircraft design and maintenance – we have very few failures of each component, not enough to determine the exact probability of each such failure, and we have even fewer situations when two components failed, so an empirical determination of such correlations is out of question.

What we do in this paper. Our objective is to use all this information to estimate the probability of failure of the entire the complex system. In this paper, we describe methods for such estimation.

2 Simplest Case: Component Failures are Independent and Failure Probabilities $P(A)$ Are Exactly Known

Let us start our analysis with the simplest case when the component failures are independent and the failure probabilities $P(A)$ for different components A are known exactly.

As we mentioned, we assume that there exists an efficient algorithm that, given a list of failed components, determines whether the whole system fails or not.

In this case, it is always possible to efficiently estimate the probability P of the system’s failure by using Monte-Carlo simulations. Specifically, we select the number of simulations N . Then, for each component A , we simulate a Boolean

variable $failing(A)$ which is true with probability $P(A)$ and false with the remaining probability $1 - P(A)$. This can be done, e.g., if we take the result r of a standard random number generator that generates values uniformly distributed on the interval $[0, 1]$ and select $failing(A)$ to be true if $r \leq P(A)$ and false otherwise: then the probability of this variable to be true is exactly $P(A)$.

Then, we apply the above-mentioned algorithm to the simulated values of the variables $failing(A)$ and conclude whether for this simulation, the system fails or not. As an estimate for the probability of the system's failure, we then take the ratio $p \stackrel{\text{def}}{=} f/N$, where f is the number of simulations on which the system failed. From statistics, it is known that the mean value of this ratio is indeed the desired probability, that the standard deviation can be estimated as $\sigma = \sqrt{p \cdot (1 - p)/N} \leq 0.5/\sqrt{N}$, and that for sufficiently large N (due to the Central Limit Theorem), the distribution of the difference $P - p$ is close to normal. Thus, with probability 99.9%, the actual value P is within the three-sigma interval $[p - 3\sigma, p + 3\sigma]$.

This enables us to determine how many iterations we need to estimate the probability P with accuracy 10% (and certainty 99.9%): due to $\sigma \leq 0.5/\sqrt{N}$, to guarantee that $3\sigma \leq 0.1$, it is sufficient to select N for which $3 \cdot 0.5/\sqrt{N} \leq 0.1$, i.e., $\sqrt{N} \geq (3 \cdot 0.5)/0.1 = 15$ and $N \geq 225$. It is important to emphasize that this number of iterations is the same no matter how many components we have – and for complex systems, we usually have many thousands of components.

Similarly, to estimate this probability with accuracy 1%, we need $N = 22,500$ iterations, etc. These numbers of iterations work for all possible values P . In practical applications, the desired probability P is small, so $1 - P \approx 1$, $\sigma \approx \sqrt{P/N}$ and the number of iterations, as determined by the condition $3\sigma \leq 0.1$ or $3\sigma \leq 0.01$, is much smaller: $N \geq 900 \cdot P$ for accuracy 10% and $N \geq 90,000 \cdot P$ for accuracy 1%.

Comment. In many cases, there are also efficient analytical algorithms for computing the desired probability of the system's failure; see, e.g., [5].

3 Important Subcase of the Simplest Case: When Components are Very Reliable

In many practical applications (e.g., in important subsystems related to aircrafts), components are highly reliable, and their probabilities of failure $P(A)$ are very small. In this case, the above Monte-Carlo technique for computing the probability P of the system's failure requires a large number of simulations, because otherwise, with high probability, in all simulations, all the components will be simulated as working properly.

For example, if the probability of a component's failure is $P(A) = 10^{-3}$, then we need at least a thousand iteration to catch a case when this component fails; if $P(A) = 10^{-6}$, we need at least a million iterations, etc.

In such situations, Monte-Carlo simulations may take a lot of computation time. In some applications, e.g., on the stage of an aircraft design, it may be OK,

but in other case, e.g., on the stage of routine aircraft maintenance, the airlines want fast turnaround, and any speed up is highly welcome.

To speed up such simulations, we can use a re-scaling idea; see, e.g., [7]. Specifically, instead of using the original values $P(A)$, we use re-scaled (larger) values $\lambda \cdot P(A)$ for some $\lambda \gg 1$. The value λ is chosen in such a way that the resulting probabilities are larger and thus, require fewer simulations to come up with cases when some components fail. As a result of applying the above Monte-Carlo simulations to these new probabilities $\lambda \cdot P(A)$, we get a probability of failure $P(\lambda)$.

In this case, one can show that while the resulting probabilities $\lambda \cdot P(A)$ are still small, the probability $P(\lambda)$ depends on λ as $P(\lambda) \approx \lambda^k \cdot P$ for some positive integer k .

Thus, to find the desired value P , we repeat this procedure for two different values $\lambda_1 \neq \lambda_2$, get the two values $P(\lambda_1)$ and $P(\lambda_2)$, and then find both unknown k and P from the resulting system of two equations with two unknowns: $P(\lambda_1) \approx \lambda_1^k \cdot P$ and $P(\lambda_2) \approx \lambda_2^k \cdot P$.

To solve this system, we first divide the first equation by the second one, getting an equation $P(\lambda_1)/P(\lambda_2) \approx (\lambda_1/\lambda_2)^k$ with one unknown k , and find $k \approx \ln(P(\lambda_1)/P(\lambda_2))/(\ln(\lambda_1/\lambda_2))$. Then, once we know k , we can find P as $P \approx P(\lambda_1)/\lambda_1^k$.

4 Cases When We Know the Probabilities $P(A)$ with Uncertainty: Fuzzy Uncertainty Can Be Reduced to Interval One

In many cases, we do not know the exact probabilities $P(A)$ of the component's failure, we only know the intervals $\mathbf{P}(A)$ that contain these probabilities, or, even more generally, a fuzzy number $\mathcal{P}(A)$ that describes this probability.

In the case of intervals, different combinations of values $P(A) \in \mathbf{P}(A)$ lead, in general, to different values P . Let us denote the dependence of P on the values $P(A)$ by f : $P = f(P(A), P(B), \dots)$. In this case, we want to know the range of possible values of the desired probability P :

$$\mathbf{P} = [\underline{P}, \overline{P}] = f(\mathbf{P}(A), \mathbf{P}(B), \dots) \stackrel{\text{def}}{=}$$

$$\{f(P(A), P(B), \dots) : P(A) \in \mathbf{P}(A), P(B) \in \mathbf{P}(B), \dots\}.$$

The problem of computing such an interval is a particular case of the general problem of interval computations, i.e., a problem of computing the range of a given function $f(x_1, \dots, x_n)$ when each of the variables x_i takes value from a given interval $\mathbf{x}_i$; see, e.g., [6, 8, 12] and references therein.

When each probability is described by a fuzzy number $\mathcal{P}(A)$, i.e., by a membership function $\mu_A(P)$ that assigns, to every real number P , a degree to which this number is possible as a value of $P(A)$, we want to find the fuzzy number

$\mathcal{P}$ that describes $f(P(A), P(B), \dots)$. A natural way to define the corresponding membership function $\mu(P)$ leads to Zadeh's extension principle:

$$\mu(P) = \sup\{\min(\mu_A(P(A)), \mu_B(P(B)), \dots) : f(P(A), P(B), \dots) = P\}.$$

It is known that from the computational viewpoint, the application of this formula can be reduced to interval computations.

Specifically, for each fuzzy set with a membership function $\mu(x)$ and for each $\alpha \in (0, 1]$, we can define this set's α -cut as $\mathcal{X}(\alpha) \stackrel{\text{def}}{=} \{x : \mu(x) \geq \alpha\}$. Vice versa, if we know the α -cuts for all α , we, for each x , can reconstruct the value $\mu(x)$ as the largest value α for which $x \in \mathcal{X}(\alpha)$. Thus, to describe a fuzzy number, it is sufficient to find all its α -cuts.

It is known that when the inputs $\mu_i(x_i)$ are fuzzy numbers, and the function $y = f(x_1, \dots, x_n)$ is continuous, then for each α , the α -cut $\mathcal{Y}(\alpha)$ of y is equal to the range of possible values of $f(x_1, \dots, x_n)$ when $x_i \in \mathcal{X}_i(\alpha)$ for all i :

$$\mathcal{Y}(\alpha) = f(\mathcal{X}_1(\alpha), \dots, \mathcal{X}_n(\alpha));$$

see, e.g., [4, 9, 13, 14]. This is how processing fuzzy data is often done – by reducing to interval computations.

$$\mathbf{P} = [\underline{P}, \overline{P}] = f(\mathbf{P}(A), \mathbf{P}(B), \dots).$$

In particular, for our problem, once know the α -cuts $\mathcal{P}(A, \alpha)$ corresponding to different components A , we can find the α -cuts $\mathcal{P}(\alpha)$ corresponding to the desired probability P as the corresponding interval range:

$$\mathcal{P}(\alpha) = f(\mathcal{P}(A, \alpha), \mathcal{P}(B, \alpha), \dots).$$

So, if we know how to solve our problem under interval uncertainty, we can also solve it under fuzzy uncertainty – e.g., by repeating the above interval computations for $\alpha = 0, 0.1, \dots, 0.9, 1.0$.

In view of this reduction, in the following text, we will only consider the case of interval uncertainty.

5 Component Failures are Independent, Failure Probabilities $P(A)$ Are Known with Interval Uncertainty: Monotonicity Case

In view of the above analysis, let us now consider the case when the probabilities $P(A)$ are only known with interval uncertainty, i.e., for each component A , we only know the interval $[\underline{P}(A), \overline{P}(A)]$ that contains the actual (unknown) value $P(A)$. We still assume that failures of different components are independent events.

Let us start with the simplest subcase when the dependence of the system's failure on the failure of components is monotonic. To be precise, we assume that

if for a certain list of failed components, the system fails, it will still fail if we add one more components to the list of failed ones. In this case, the smaller the probability of failure $P(A)$ for each component A , the smaller the probability P that the system as a whole will fail. Similarly, the larger the probability of failure $P(A)$ for each component A , the larger the probability P that the system as a whole will fail.

Thus, to compute the smallest possible value $\underline{P}$ of the failure probability, it is sufficient to consider the values $\underline{P}(A)$. Similarly, to compute the largest possible value $\overline{P}$ of the failure probability, it is sufficient to consider the values $\overline{P}(A)$. Thus, in the monotonic case, to compute the range $[\underline{P}, \overline{P}]$ of possible values of overall failure probability under interval uncertainty, it is sufficient to solve two problems in each of which we know probabilities with certainty:

- to compute $\underline{P}$, we assume that for each component A , the failure probability is equal to $\underline{P}(A)$;
- to compute $\overline{P}$, we assume that for each component A , the failure probability is equal to $\overline{P}(A)$.

6 In Practice, the Dependence is Sometimes Non-Monotonic

Let us show that in some practically reasonable situations, the dependence of the system's failure on the failure of components is non-monotonic. This may sound counter-intuitive at first glance: adding one more failing component to the list of failed ones suddenly makes the previously failing system recover, but here is an example when exactly this seemingly counter-intuitive behavior makes perfect sense.

To increase reliability, systems include duplication: for many important functions, there is a duplicate subsystem ready to take charge if the main subsystem fails. How do we detect that the main system failed? Usually, a subsystem contains several sensors; sensors sometimes fail, as a result of which their signal no longer reflect the actual value of the quantity they are supposed to measure. For example, a temperature sensor which is supposed to generate a signal proportional to the temperature, if failed, produces no signal at all, which the system will naturally interpret as a 0 temperature. To detect the sensor failure, subsystems often use statistical criteria. For example, for each sensor i , we usually know the mean m_i and the standard deviation σ_i of the corresponding quantity. When these quantities are independent and approximately normally distributed, then, for the measurement values x_i , the sum $X^2 \stackrel{\text{def}}{=} \sum_{i=1}^n \frac{(x_i - m_i)^2}{\sigma_i^2}$ is the sum of n standard normal distributions and thus, follows the chi-square distributed with n degrees of freedom. So, if the actual value of this sum exceeds the threshold corresponding to confidence level $p = 0.05$, this means that we can confidently conclude that some of the sensors are malfunctioning.

If the number n of sensors is large, then one malfunctioning sensor may not increase the sum X^2 too high, and so, its malfunctioning will not be detected,

and the system will fail. On the other hand, if all n sensors fail, e.g., show 0 instead of the correct temperature, each term in the sum will be large, the sum will exceed the threshold – and the system will detect the malfunctioning. In this case, the second redundant subsystem will be activated, and the system as a whole will thus continue to function normally.

This is exactly the case of non-monotonicity: when only one sensor fails, the system as a whole fails; however, if, in addition to the originally failed sensor, many other sensors fail, the system as a whole becomes functioning well.

7 What If We Have Few Components, With Respect to Which the Dependence is Non-Monotonic

Let us start with the simplest case when there are only few components with respect to which the dependence is non-monotonic, and with respect to all other components, the dependence *is* still monotonic. In this case, for all monotonic components B , as before, we can take $P(B) = \bar{P}(B)$ when we are computing $\bar{P}$, and take $P(B) = \underline{P}(B)$ when we are computing $\underline{P}$.

For non-monotonic components A , for computing each of the values $\underline{P}$ and $\bar{P}$, we need to take into account all possible values $P(A) \in [\underline{P}(A), \bar{P}(A)]$. For each such component, by using the formula of full probability, we can represent the probability P of the system's failure as follows:

$$P = P(A) \cdot P(F|A) + (1 - P(A)) \cdot P(F|\neg A),$$

where $P(F|A)$ is the conditional probability that the system fails under the condition that the component A fails, and $P(F|\neg A)$ is the conditional probability that the system fails under the condition that the component A does not fail. The conditional probabilities $P(F|A)$ and $P(F|\neg A)$ do not depend on $P(A)$, so the resulting dependence of P on $P(A)$ is linear. A linear function attains its minimum and maximum at the endpoints. Thus, to find $\underline{P}$ and $\bar{P}$, it is not necessary to consider all possible values $P(A) \in [\underline{P}(A), \bar{P}(A)]$, it is sufficient to only consider two values: $P(A) = \underline{P}(A)$ and $P(A) = \bar{P}(A)$.

For each of these two values, for another non-monotonic component A' , we have two possible options $P(A') = \underline{P}(A')$ and $P(A') = \bar{P}(A')$; thus, in this case, we need to consider $2 \times 2 = 4$ possible combinations of values $P(A)$ and $P(A')$.

In general, when we have k non-monotonic components $A_1, \dots, A_k$, it is sufficient to consider 2^k possible combinations of values $\underline{P}(A_i)$ and $\bar{P}(A_i)$ corresponding to each of these components. When k is small, this is doable – and, as our preliminary experiments show, works very well.

This procedure requires times which grows as 2^k . As we mentioned earlier, when k is large, the needed computation time becomes unrealistically large.

8 General Case: the Problem is NP-Hard, and Even Checking Whether a Component Is Monotonic Is NP-hard

Natural question. The fact that the above algorithm requires unrealistic exponential time raises a natural question: is it because our algorithm is inefficient or is it because the problem itself is difficult?

The problem is NP-hard. In the general case, when no assumption is made about monotonicity, the problem is as follows:

- we have a propositional formula F with n variables A_i – each variable A_i is true if the corresponding component fails, and the formula F is true if the system as whole fails;
- for each component i , we know the interval $[P(A_i), \overline{P}(A_i)]$ that contains the actual (unknown) $P(A_i)$ that this component fails;
- we assume that the failures of different components are independent events.

Different values $P(A_i) \in [P(A_i), \overline{P}(A_i)]$ lead, in general, to different values of the probability P that F is true (i.e., that the system failed). Our objective is to compute the range $[P, \overline{P}]$ of possible values of this probability.

One can easily show that, in general, this problem is NP-hard (for precise definitions, see, e.g., [15]). Indeed, it is well known that the following *propositional satisfiability* problem SAT is NP-hard: given a propositional formula F , check whether this formula is satisfiable, i.e., whether there exist values A_i that make it true. We will prove that our problem is NP-hard by reducing SAT to it: for every particular case F of SAT, there is a particular case of our problem whose solution leads to a solution to the F . Indeed, for every formula F , let us take $[P(A_i), \overline{P}(A_i)] = [0, 1]$ for all variables i .

If the formula F is not satisfiable, then F is always false, so the probability of its being true is 0. In this case, the range of possible values of P consists of a single value 0: $[P, \overline{P}] = [0, 0]$.

On the other hand, if F is satisfiable, e.g., if F is true for A_1 true, A_2 false, etc., then we can take $P(A_1) = 1$, $P(A_2) = 0$, etc., and conclude that under these probabilities, F is always true, i.e., $P = 1$. Thus, in this case, $\overline{P} = 1$.

So, by computing $\overline{P}$, we will get either $\overline{P} = 0$ or $\overline{P} = 1$. In the first case, F is satisfiable, in the second case, it is not. The reduction proves that our problem is NP-hard.

Even checking monotonicity is NP-hard. Let us show that even checking monotonicity is NP-hard. Intuitively, monotonicity with respect to a component A means that if the system was in a failing state, and we change the state of A to failing, the system remains failing. In precise terms, monotonicity means that if the formula F was true, and we change the value of the variable A from false to true, then F remains true.

Let us prove that the problem of checking monotonicity is NP-hard by reducing SAT to this problem. Indeed, for every propositional formula $F(A_1, \dots, A_n)$, we can form a new formula $F' \stackrel{\text{def}}{=} F(A_1, \dots, A_n) \& \neg A_0$.

When the original formula F is not satisfiable, then F is always false and thus, the new formula F' is also always false. In this case, F' is, by definition, monotonic with respect to A_0 .

When F is satisfiable, this means that F is equal to “true” for some values of $A_1, \dots, A_n$. In this case, for $A_0 = \text{“false”}$, the new formula F' is true, but if we make $A_0 = \text{“true”}$, F' becomes false. Thus, in this case, F' is not monotonic with respect to A_0 .

Summarizing: the new formula F' is monotonic with respect to A_0 if and only if the original formula F was satisfiable. This reduction proves that checking monotonicity is indeed NP-hard.

9 Case of Narrow Intervals: Cauchy Deviate Method

In many practical situations, intervals are narrow. In this case, we can use an efficient Cauchy deviates method to find the range of the resulting probability P ; see, e.g., [10].

10 What If We Do Not Assume Independence

Exact methods. If we do not assume independence, then, in principle, we can have different probabilities P of failure even when the failure probabilities $P(A_i)$ of all components are known exactly. For example, if the system consists of two components and it fails if both components fail, i.e., if $F = A_1 \& A_2$, then possible values of P take an interval

$$[\underline{P}, \overline{P}] = [\max(P(A_1) + P(A_2) - 1, 0), \min(P(A_1), P(A_2))].$$

In general, to describe probabilities of all possible combinations of statements A_i , it is sufficient to describe 2^n probabilities of *atomic* statements $A_1^{\varepsilon_1} \& \dots \& A_n^{\varepsilon_n}$, where $\varepsilon_i \in \{-, +\}$, A^+ means A , and A^- means $\neg A$. These probabilities satisfy the condition that their sum is 1; each given probability $P(A_i)$ and the desired probability P can be described as a sum of some such probabilities, so the problem of finding the range of P becomes the particular case of *linear programming* problems, when we need to find the minimum and maximum of a linear function under linear constraints; see, e.g., [16].

For example, in the above case $n = 2$, we have four non-negative unknowns $P_{++} = P(A_1 \& A_2)$, $P_{+-} = P(A_1 \& \neg A_2)$, $P_{-+} = P(\neg A_1 \& A_2)$, and $P_{--} = P(\neg A_1 \& \neg A_2)$ that satisfy the constraints $P_{++} + P_{+-} + P_{-+} + P_{--} = 1$, $P_{++} + P_{+-} = P(A_1)$, and $P_{++} + P_{-+} = P(A_2)$. Here, $P = P_{++}$, so, e.g., to find $\underline{P}$, we need to minimize P_{++} under these constraints.

Limitations of the exact methods. This works well if n is small, but for large n , this method requires an unrealistically long time.

Heuristic approximate methods. In principle, we can use technique similar to straightforward interval computations. Indeed, for simple formulas F like $\neg A_1$, $A_1 \vee A_2$, or $A_1 \& A_2$, we have explicit formulas for the range of the probability $P(F)$. So, to estimate the range of the probability P for an arbitrary formula F , we can do the following:

- we *parse* the expression F , i.e., represent it as a sequence of simple boolean operations – the same sequence that a computer computing F would follow, and
- replace each computation step with corresponding operations with probability ranges.

In this algorithm, at each moment of time, we keep the bounds on the probabilities $P(A_i)$ and on the probabilities $P(F_j)$ of the corresponding intermediate formulas.

The problem with this approach is that at each step, we ignore the dependence between the intermediate results F_j ; hence intervals grow too wide. For example, the estimate for $P(A \vee \neg A)$ computed this way is not 1, but an interval containing the correct value 1.

A more accurate algorithm was proposed in [1–3]. In this algorithm, at each stage, besides the bounds on the values $P(F_j)$ (including the original bounds on $P(A_i)$, or – if available – exact values of $P(A_i)$), we also compute the bounds for the probabilities $P(F_j \& F_k)$, $P(F_j \& \neg F_k)$, $P(\neg F_j \& F_k)$, and $P(\neg F_j \& \neg F_k)$.

On each computation step, when we add a new intermediate result F_a (e.g., $F_a = F_b \& F_c$), we add bounds for the probabilities of the new statement F_a and of all possible combinations that include F_a , i.e., on the probabilities $P(F_a \& F_k)$, $P(F_a \& \neg F_k)$, $P(\neg F_a \& F_k)$, and $P(\neg F_a \& \neg F_k)$. To compute each of these probabilities, we use the known bounds on the probabilities of combinations of F_b and F_k , F_c and F_k , F_b and F_c , and get the desired bounds on the combinations of F_a and F_k by solving the corresponding linear programming problem. In this linear programming problem, we consider, as variables, $2^3 = 8$ probabilities of atomic statements $F_b^{\varepsilon_b} \& F_c^{\varepsilon_c} \& F_k^{\varepsilon_k}$.

At the end of the process, we get an interval $\tilde{P}$. Similarly to interval computations, we can prove, by induction, that every possible value P of the system's failure is contained in this interval, i.e., that the interval $\tilde{P}$ is an enclosure for the desired range $[\underline{P}, \overline{P}]$: $[\underline{P}, \overline{P}] \subseteq \tilde{P}$.

This algorithm requires more computation time than the above straightforward algorithm – since for s intermediate steps, we now need s^2 estimations instead of s – but as a result, we get more accurate results, with smaller “excess width” for the resulting enclosure. For example, the probability $P(A \vee \neg A)$ is estimated as 1.

To get an even better accuracy, instead of probabilities of all possible combinations of two intermediate results, we can compute, on each step, probabilities of all possible combinations of three such results: F_i , F_j , and F_k . In this case, computation time grows as s^3 but the resulting enclosure is even more accurate. Similarly, we can have combinations of fours, fives, etc.

Acknowledgments. C. Jacob was supported by a grant from @MOST Prototype, a joint project of Airbus, Institut de Recherche en Informatique de Toulouse (IRIT), Laboratoire d'Analyse et d'Architecture des Systèmes (LAAS), ONERA, and Institut Supérieur de l'Aéronautique et de l'Espace (ISAE).

V. Kreinovich was supported in part by the National Science Foundation grants HRD-0734825 and DUE-0926721 and by Grant 1 T36 GM078000-01 from the National Institutes of Health.

References

1. Ceberio, M., Kreinovich, V., Chopra, S., Longpré, L., Nguyen, H. T., Ludaescher, B., and Baral, C.: Interval-type and affine arithmetic-type techniques for handling uncertainty in expert systems, *Journal of Computational and Applied Mathematics*, 199(2), 403–410 (2007)
2. Chopra, S.: Affine arithmetic-type techniques for handling uncertainty in expert systems, Master's thesis, Department of Computer Science, University of Texas at El Paso, 2005.
3. Chopra, S.: Affine arithmetic-type techniques for handling uncertainty in expert systems, *International Journal of Intelligent Technologies and Applied Statistics*, 1(1), 59–110 (2008)
4. Dubois, D., Prade, H.: Operations on fuzzy numbers, *International Journal of Systems Science*, 9, 613–626 (1978)
5. Dutuit, Y., Rauzy, A.: Approximate estimation of system reliability via fault trees, *Reliability Engineering and System Safety*, 87(2), 163–172 (2005)
6. Interval computations website <http://www.cs.utep.edu/interval-comp>
7. Jaksurat, P., Freudenthal, E., Ceberio, M., Kreinovich, V.: Probabilistic Approach to Trust: Ideas, Algorithms, and Simulations, *Proceedings of the Fifth International Conference on Intelligent Technologies InTech'04*, Houston, Texas, December 2–4, 2004 (2004)
8. Jaulin, L., Kieffer, M., Didrit, O., Walter, E.: *Applied Interval Analysis*, Springer, London (2001)
9. Klir, G., Yuan, B.: *Fuzzy Sets and Fuzzy Logic*, Prentice Hall, Upper Saddle River, NJ (1995)
10. Kreinovich, V., Ferson, S.: A new Cauchy-based black-box technique for uncertainty in risk analysis, *Reliability Engineering and Systems Safety*, 85(1–3), 267–279 (2004)
11. Kreinovich, V., et al.: *Computational Complexity and Feasibility of Data Processing and Interval Computations*, Kluwer, Dordrecht (1997)
12. Moore, R. E., Kearfott, R. B., Cloud, M. J.: *Introduction to Interval Analysis*, SIAM Press, Philadelphia, Pennsylvania (2009)
13. Nguyen, H. T., Kreinovich, V.: Nested intervals and sets: concepts, relations to fuzzy sets, and applications, In: Kearfott, R. B., Kreinovich, V., eds., *Applications of Interval Computations*, Kluwer, Dordrecht, 245–290 (1996)
14. Nguyen, H. T., Walker, E. A.: *A First Course in Fuzzy Logic*, Chapman & Hall/CRC, Boca Raton, Florida (2006)
15. Papadimitriou, C.: *Computational Complexity*, Addison Welsey, Reading, Massachusetts (1994)
16. Walley, P.: *Statistical reasoning with imprecise probabilities*. Chapman & Hall: New York (1991)