

Risk_OMT – Hybrid approach

Author: Jørn Vatn, NTNU, Department of Production and Quality Engineering

Participants of PK8200

Document created: Spring 2011

Last update: 2013-01-05

1 Introduction

This document describes the technical aspects of the hybrid implementation of the Risk_OMT model. The Risk_OMT offers two modelling approaches. The full-BBN (Bayesian Belief Network) approach utilize BBN modelling both for the soft influences between risk influencing factors (RIFs) and the formal probabilistic relations described by fault- and event trees. The hybrid implementation of the Risk_OMT model uses a BBN specification of the relation between the RIFs but uses ordinary processing of the fault- and event trees.

In the fault and event trees failure of an activity is divided into *failures of omission* and *failure of execution*. Failure of omission denotes whether or not the prescribed activity is carried out. Failure of execution denotes inadequate actions that may cause failures, e.g., acts performed in a wrong sequence, at wrong time, without required precision etc. Failure of execution is seen as results of *human errors* and *violations*. Human error is further divided into *mistakes* and *slips & lapses*, where mistakes involve actions that are based on failure of interpretation of procedures, and/or failures of judgemental/inferential processes involved in the prescribed activity. This category does not distinguish between whether or not the actions directed by this judgement activities run according to the actor's plan. Typical mistakes are inadequate judgement/conclusion due to intrinsic conditions such as competence, fatigue, mode etc, and extrinsic conditions such as communication, information, work load, time pressure etc. Slips & lapses involve actions that represent unintended deviation from those practiced represented in the formal procedures. This is deviation due to error in execution and/or the storage stage of an action sequence. For our purpose, this category represents only actions where there is no intended violation, failure of interpretation of procedures and judgement failures prior to the action carried out. In the Risk_OMT model separate BBNs are developed for the RIF structure for (i) mistakes, (ii) slips & lapses, and (iii) violations. For failures of omission currently no RIF model is derived.

2 Risk influencing factors

A Risk Influencing Factor (RIF) represents a condition or a situation that influences the risk in a risk model. In this presentation we always assume that the RIFs are influencing the risk through parameters used in the risk model. In the Risk_OMT we mainly focus on different organisational conditions that have a theoretical and/or empirical grounded influence on the possible deviations from required actions, and hence should be reflected in probability assignment of errors or failures. Further, Risk_OMT operates with 2 levels of RIFs which links the organisational conditions (RIF Level 1) to strategic management decisions (RIF Level 2). In the current Risk_OMT implementation it is only RIFs on level 1 that directly influence the basic event probabilities. Note that the term risk is interpreted differently within the society of risk analysis. In a classical risk analysis framework risk is seen as a property of the system being analysed. Further

probabilities in a risk model are considered to represent some true likelihood of e.g., component failures and human errors. With such an interpretation we may think of the RIFs as a way to establish a true causal link between some conditions and the basic event probabilities. In an epistemic interpretation of risk the main focus is on uncertainty. Risk is essential uncertainty regarding the occurrence and severity of undesired events. In such a framework basic event probabilities are not considered as some true values, but are expressions of our uncertainty regarding the occurrence of the basic events. The RIFs will then represent conditions that we take into account when assigning probabilities (expressing uncertainty) to the basic events, but we do not consider any causal link as for the classical interpretation.

The Risk_OMT modelling framework is an extension of the BORA release model (Aven et.al. 2006). There are two major changes in the Risk_OMT model compared to the BORA release model. Whereas the BORA release model combined the RIFs on the same level, the Risk_OMT model introduces a hierarchy between the RIFs. Further the BORA release model considered the RIFs to be known without any uncertainty. In the Risk_OMT model RIFs are still considered to be theoretical constructs that influences the risk, but we do not have exact knowledge regarding the value of the RIFs, and hence they are treated as stochastic variables (random quantities).

Formally we use the term *score* to denote the summarized information regarding the RIFs from interviews, surveys etc. A score is thus treated as a realization (observation) of the true underlying RIF. In the BBN this corresponds to an arrow from the RIF to the corresponding score. The scoring system is based on characters A to F, where A corresponds to best industry practice, and F corresponds to an unacceptable state with respect to the actual RIF.

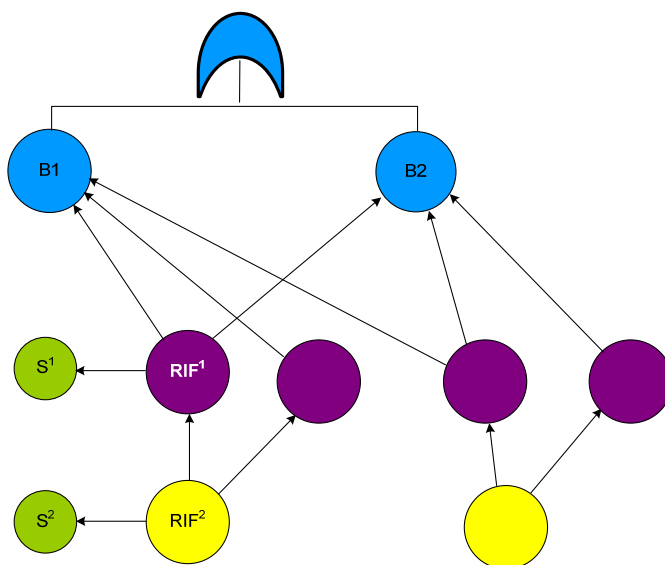


Figure 1 RIFs on two levels with scores and relation to basic events in a fault tree

Figure 1 shows an example of a RIF structure. Level 1 RIFs point to the basic events in the fault tree showing that level 1 RIFs influences the basic events. Level 2 RIFs influence the level 1 RIFs, and there is an arrow from the RIFs to the scores to indicate that the scores are treated as realizations (observations) of the true underlying RIFs.

Examples of level 1 RIFs are *technical documentation* and *time pressure*. Corresponding level 2 RIFs are *management of information* and *management of tasks* respectively.

3 Impact of the level one RIFs on the basic events

In the hybrid Risk_OMT model the impact of the RIFs are explicitly modelled via the probability of occurrence of a basic event or a barrier in the fault or event tree. We now consider basic event number i . Three quantities span the sample space for the basic event probability for this event:

$q_{i,A}$ = average basic event probability corresponding to average industry practice, i.e., all RIFs equal to the character C.

$q_{i,L}$ = lowest basic event probability corresponding to the best practice in the industry, i.e. all RIFs equal to the character A.

$q_{i,H}$ = highest basic event probability corresponding to the an unacceptable industry practice, i.e. all RIFs equal to the character F. It is not expected to observe RIFs of character F.

Often $q_{i,L}$ and $q_{i,H}$ are specified indirectly by error factors $EF_{i,H} = q_{i,H} / q_{i,A}$ and $EF_{i,L} = q_{i,A} / q_{i,L}$. In the modelling it will be convenient to represent the RIFs by numeric values rather than character values. It is convenient to map the character values on the interval [0,1]. Splitting this interval into 6 sub intervals, and mapping the character value into the centre gives the value 1/12 for an A, the value 1/6+1/12 = 3/12 for a B , the value 2/6 + 1/12 = 5/12 for a C up to 11/12 for a value F. etc. In the following we will always use this mapping in the numeric quantifications. In some situations we may use a more differentiated labelling than the pure characters, i.e., we may succeed the characters with extra plusses (+) or minuses (-). For example we may use that A++ corresponds to a value 0, A+ corresponds to a value 1/24 etc. If we use r as a value of a weighted sum of the RIFs influencing the basic event probability we now introduce $q_i(r)$ to describe the functional relationship between the RIF value (r) and the basic event probability. We have that $q_i(0) = q_{i,L}$, $q_i(5/12) = q_{i,A}$ and $q_i(1) = q_{i,H}$. In between these values we may either use linear or geometric interpolation. For high error factors it is recommended to use a geometric interpolation.

We will now assume that there are totally J RIFs that are influencing basic event i . Let $\mathbf{R} = [R_1, R_2, \dots, R_J]$ be a vector of stochastic variables to represent these (standardized) RIFs, and let $p_{\mathbf{R}}(\mathbf{r}) = \Pr(R_1 = r_1, R_2 = r_2, \dots, R_J = r_J)$ be the joint probability distribution over these RIFs. Each RIF might have different weight with respect to the influence on the basic event probability. Now let w_j be standardized weight for RIF j . A first approximation for the total impact of the RIFs on the basic event probability is given by:

$$q_i = \Pr(\text{Failure of basic event } i) = \sum_{\mathbf{r}} q_i(\sum_j w_j r_j) p_{\mathbf{R}}(\mathbf{r}) \quad (1)$$

where $\sum_{\mathbf{r}}$ represents the sum over all possible values of $\mathbf{r}$. Equation (1) is then used to establish the basic event probabilities to use in the fault and event tree part of the hybrid risk analysis.

4 The beta distribution to describe uncertainty regarding the RIFs

A mathematical convenient probability distribution to use for continuous variables on the interval [0,1] is the beta distribution. Although the scoring of the RIFs are on an ordinal level, a continuous ratio scale seems appropriate for the modelling. The probability density function of the beta distribution is given by:

$$f(r) = r^{\alpha-1} (1-r)^{\beta-1} / B(\alpha, \beta) \quad (2)$$

where $B(\alpha, \beta) = \Gamma(\alpha)\Gamma(\beta)/\Gamma(\alpha + \beta)$ is the beta function, and $\Gamma(\cdot)$ is the gamma function. α and β are parameters in the distribution.

If R is beta distributed with parameters α and β the expected value and variance are given by:

$$E(R) = \frac{\alpha}{\alpha + \beta} \quad (3)$$

$$\text{Var}(R) = \frac{\alpha\beta}{(\alpha + \beta)^2(\alpha + \beta + 1)} \quad (4)$$

Note that the beta distribution represents a *conjugate* prior distribution for the binomial distribution. Thus if the beta distribution is used to describe the parameter r in a binomial distribution with prior parameters α_0 and β_0 , then the posterior distribution is also beta distributed with parameters $\alpha_0 + x$ and $\beta_0 + n - x$ where x is the number of successes and n is the number of trials of an experiment provided to update the prior distribution, i.e., the posterior distribution is a beta distribution with parameters

$$\alpha = \alpha_0 + x \quad (5)$$

$$\beta = \beta_0 + n - x \quad (6)$$

5 Updating the RIF distributions based on the scores

The above results in equations (5) and (6) do not apply directly to our situation since we will not get observations from a binomial trial but rather one observation considered to be a realisation of the true RIF. Let α_0 og β_0 be the parameters in the prior distribution of the RIF prior to observing the score S . Given the true value of the RIF, say r , it is reasonable to assume that $E(S|r) = r$, and further we assume that it is possible to specify $\text{Var}(S|r) = V_S$. We now make the following argument: We will use the value of the score, say s , by translating the information to a binomial situation, e.g., finding x and n . This is done due to the simple result that exists for the binomial situation. Since X/n is an estimator for r in the binomial situation, and the score S is the estimator for r in our situation, it seems reasonable to require:

$$\text{Var}(X/n) = r(1 - r)/n = \text{Var}(S) = V_S \quad (7)$$

Thus, if we know V and replace r with it's estimate s , we should have:

$$n = s(1 - s)/\text{Var}(S) = s(1 - s)/V_S \quad (8)$$

further since x/n and s both are estimates of r , we set $x = s \cdot n$. Utilizing the result from the binomial situation where the posterior distribution is beta distributed with parameters $\alpha_0 + x$ and $\beta_0 + n - x$ we will in our situation approximate the posterior distribution with a beta distribution with parameters:

$$\alpha = \alpha_0 + s^2(1-s)/V_S \quad (9)$$

$$\beta = \beta_0 + s(1-s)/V_S - s^2(1-s)/V_S = \beta_0 + s(1-s)^2/V_S \quad (10)$$

Exercise 1

Find the expected value and the variance of the posterior distribution with the parameters obtained by equations (9) and (10). Compare this result with the expected value and variance of the weighted sum of the prior mean and the score where the reciprocal variances are used as weights.

5.1 Level 2 RIFs

For level 2 RIFs it is straight forward to use the result in equations (9) and (10) to find posterior distributions for the RIFs. Various principles may be used for specifying the prior distribution. In order to have a method that is data driven as far as possible, it seems reasonable to apply $\alpha_0 = \beta_0 = 0.5$ corresponding to Jeffreys prior (Jeffreys, 1946).

Exercise 2

Apply Jeffreys prior together with equations (9) and (10) in order to find the expected value and the variance of the posterior distribution for scores corresponding to the characters A, B, ..., F. Present the result in a table for V_s equal to 0.2^2 , 0.1^2 and 0.05^2 respectively.

5.2 Level 1 RIFs

We will start by assigning the posterior distribution of level 1 RIFs given the value of the parent RIF, i.e., the corresponding level 2 RIF denoted P (i.e., parent). Given the value of the level 2 RIF, say $P=p$, it is reasonable to specify a prior distribution of the RIF, say R , with expected value $E(R|P=p) = p$. Thus the structural dependencies between the parent RIF and the child RIF is considered to give the same expectation. But how strong is the structural dependency, i.e., the influence of the parent RIF on the child RIF. Such a structural dependency may be expressed by the variance of the child RIF, i.e., $\text{Var}(R|P=p)$. To make the model simple we assume that $\text{Var}(R|P=p) = \text{Var}(R) = V_p$ where it is possible to specify V_p independent of the actual value of the parent. V_p may be considered as a measure of the *structural dependence* in the model. Proposed values for the structural dependency are $V_p = 0.2^2$ (low dependency), $V_p = 0.1^2$ (medium dependency) and $V_p = 0.05^2$ (high dependency). Prior to observing the score, it seems reasonable to express the prior distribution of the RIF with a beta distribution with expected value p and variance V_p .

Exercise 3

Use equations (3) and (4) to show that we may obtain the prior parameters in this situation by:

$$\beta_0 = \left(\frac{p(1-p)}{V_p} - 1 \right) (1 - p) \quad (11)$$

$$\alpha_0 = \frac{p\beta_0}{(1-p)} \quad (12)$$

Conditional on the value of the parent level 2 RIF; i.e., $P=p$, and the structural influence between these RIFs, the prior distribution may be obtained by applying the parameters in equations (11) and (12). Given the score $S=s$ of the level 1 RIF we apply equations (9) and (10) to find the conditional posterior distribution, i.e., given the parent value. In order to find the unconditional posterior distribution, we may integrate over the posterior distribution of the parent node.

It is, however, important to stress that we do not need the unconditional posterior distribution of the child RIFs, i.e., the level 1 RIFs. In equation (1) we need the joint distribution over the level 1 RIFs that directly influences the basic event probability. From the theory of BBN, we know that the level 1 RIFs are independent *given* their parents, i.e, the level 2 RIFs. This means that we may multiply the conditional posterior distributions for level 1 RIFs to find the required $p(\mathbf{r})$ in equation (1) and then integrate over the joint posterior distribution of the level 2 RIFs.

In equation (1) we did assume that the RIFs were made discrete, i.e., each RIF takes a finite number of values. This is done in order to simplify calculations. Since the scores are measured on six different values, it seems reasonable to use 6 values for each RIF both on level 1 and level 2. Then we may for the posterior distribution of the level 2 RIFs calculate a point probability for each interval, i.e., $[0,1/6]$, $[1/6,2/6]$ etc. Similar, given the (discrete) values of the level 2 RIFs, we may calculate point probabilities for the level 1 RIFs, and $p(\mathbf{r})$ is then found by applying the law of total probability.

Exercise 4

Write a simple code (visual basic, matlab, fortran or C) to find the point probabilities for each interval $[0,1/6]$, $[1/6,2/6]$, ..., given the parameters in the posterior distribution.

Exercise 5

Consider a situation with one level 2 RIF, two level 1 RIFs influenced by the level 2 RIF. Write a simple code to find the unconditional distribution over the weighted sum of the level 1 RIFs. Make the code flexible such it is possible to specify the weights, the scores, the structural influences V_p 's and the variances of the scores V_s 's.

Exercise 6

Discuss extension of the model in the previous exercise where there are more than two level 1 RIFs for each level 2 RIF, and where there are more than one level 2 RIF. Hint: Since each level 1 RIF is influenced by one and only one level 2 RIF, the subset of level 1 RIFs with common parent level 2 RIF may be treated separately. Discuss why this will reduce the number of combinations to run through. Also discuss how to implement the solution if this savings should be obtained.

Exercise 7

Assume that you have n RIFs where each RIF may take m different values. Propose an algorithm to generate all possible combinations of the n RIFs.

6 Interactions between RIFs

The influence of one RIF on the basic event probabilities is assumed to be independent of the value of the other RIFs in the basic Risk_OMT model. In many situations it might be reasonable to believe that for example the negative influence of a very bad RIF is higher if the one or more of the other RIFs also have a very bad value compared to more moderate values of these RIFs. It might also be argued that the positive influence of a very good RIF is higher if one or more of the other RIFs are good, compared to these RIFs having average or bad values. Finally, there might be cases where a good value of one or more RIFs

balances or neutralizes the bad values of other RIFs. In this study we will only treat negative effects where bad values of two or more RIFs strengthened the negative influence on the basic event.

The arguments in the modelling of interaction effects are as follows. Interaction effects are only modelled between the level 1 RIFs due to the fact that no level 1 RIFs are influenced by more than one level 2 RIFs in the current model, hence it make no sense to introduce interaction effects between the level 2 RIFs.

Further the influence of the level 1 RIFs on the basic event probabilities are essentially in the basic Risk_OMT model determined by a weighted sum of the Level 1 RIFs, i.e.,

$$r = \sum_i w_i r_i \quad (13)$$

where r_i is a value of RIF number i on level 1, and where we assume that the RIFs values are measured in numerical units and not by letters in the quantification part. In the modelling the RIFs are treated as random variables, hence we need to integrate over equation (13) over the simultaneous distribution to the RIFs. But as is seen from equation (13) no interaction effects are introduced. In the modelling of interaction effects sub sets of the total set of RIFs are uuconsidered to represent a potential for interaction. However, we have assumed that interaction effects only come into play when all the RIFs in a subset have a value worse than the average value, i.e., a C. In the Risk_OMT project we only consider sub sets of two or three RIFs. In order to simplify the modelling of interaction effects we assign a weight, w_l of the interaction effect which is relative to the weight of the various RIFs in the interaction sub set, say l . For each RIF in the sub set l we then may find a total weight of the RIF in addition the original weight of the RIF, i.e.,

$$w_{l,i} = w_i w_l f \quad (14)$$

where w_i is the original weight of the RIF, and f is a correction factor. If one or more of the RIFs have a value better than the average RIF-value (C) we set $f = 0$. If all the RIF values in l have the worst value (F), we set $f = 1$. For values between we apply a linear transformation:

$$f = (\sum r - \sum C) / (\sum F - \sum C) \quad (15)$$

where $\sum r$ is the sum of RIF values in l , $\sum C$ is the sum of the same RIFs if they all equal "C", and $\sum F$ is the sum if they all equal "F". It is then easy to verify that if all RIFs take the value "C" we get $f = 0$, and if all RIFs take the value "F" we get $f = 1$ corresponding to our assumptions. The total impact of the RIFs on the basic event probability is now found:

$$r = \sum_i w_i r_i + \sum_{i \in l} w_{l,i} r_i \quad (16)$$

where we have summed the interaction effects for one sub set of interactions. In principle there might be more than one sub set of interaction effects, and we need to add these also.

Since we have added more terms in equation (16) we might consider to reduce the original weights to prevent the weighted sum to exceed the worst possible value (corresponding to an F). But since we do not open for a similar "positive" interaction effects, such a change in the RIFs will also result in a situation where only RIF values equal to the average (C) will mot result in the average failure probability in the fault tree. Therefore, we rather accept that the model becomes more conservative for very bad values of the RIFs included in the sub set of interactions, i.e., we do not adjust the weights.

7 Common cause modelling

Podofolini et. al. (2009) have reviewed decision tree models for assessing human reliability analysis dependency. They refer to the common practice to introduce dependence levels, say zero, low, moderate, high and complete. For each of these levels it corresponds a β -factor which is the conditional probability of a subsequent failure given a first failure. Further they report that the literature fails to be consistent with respect to which factors affect the dependence levels, and to which strength. Podofolini et. al. (2009) summarizes the literature both with respect to factors included, and their importance. Note that common cause failures seem to be a more important issue *after* a critical event compared to our situation where we are modelling what is happening prior to the leak. Therefore we will assume lower common cause influence than found in the literature. The following factors are considered most important, see Podofolini et. al. (2009) for further elaboration and discussion related to the literature:

- Closeness in time
- Similarity of crew/performer(s)
- Stress
- Complexity

To simplify the modelling we introduce a scoring regime for each of the factors. Let S_i denote the score of the i^{th} factor. It then seems reasonable to introduce a relation between the β -factor and the scores:

$$\beta = \beta_0 \prod_i w_i^{S_i} \quad (17)$$

Where β_0 is a baseline common cause factor, and w_i is the weight of factor i . Table 1 shows the principles for setting scores, and the weights used for each of the factors.

Table 1 Weights and principles for setting scores in the CCF model

Dependency factor	Weight	Scores	
		Best = -1	Worst = 1
Closeness in time	2	The closeness in time is assumed to depend on the type of tasks considered. The following scores are proposed for the relevant situations. Control planning Planning, $S = 1$ Control Execution, $S = -1$ Execution Execution, $S = \frac{1}{2}$ Note, we assume there are no dependencies between planning activities and execution activities.	
Similarity of crew	3	Different crew, $S = -1$	Same crew, $S = 1$
Stress	2	The stress level is based on the RIF for <i>time pressure</i>. Let r denote the linear mapping of the RIF on the interval (0,1) where A corresponds to 0, and F corresponds to 1. The score of the stress dependency factor is then given by $S = 2(r - \frac{1}{2})$	
Complexity	1.5	The Risk_OMT model does not include any RIF explicitly used to describe complexity. The RIF for <i>design</i> and <i>HMI</i> are considered to be the most relevant RIFs indicating complexity. If the values of these are denoted r_1 and r_2 respectively, the score of the complexity dependency factor is then given by $S = (r_1 + r_2 - 1)$	

The baseline dependency level is set to $\beta_0 = 0.05$ for failure types “violation”, “omission”, and “mistake”. For “slips & lapses” the common cause problem is considered slightly lower, and the value $\beta_0 = 0.03$ is used.

There are two feasible ways to include the common cause effects in the modelling. One way is to model explicitly the common cause effects by introducing additional basic events in the fault and event trees. For a full BBN model this is the only way to represent such common cause effects. If the hybrid approach with a mixture of BBN models and event and fault trees is used, we may also introduce common cause failures in the post-processing of the minimal cut sets. The challenge then is to describe the possible dependencies for various classes of basic events, and then add common cause terms when the minimal cut set contributions are calculated. For example if a minimal cut set comprises the following basic events: {P=Planning error, CP=Control Planning error, E=Execution error, CE=Control Execution error} and we introduce $\beta_{CP|P}$ and $\beta_{CE|E}$ as common cause factors for controlling the plan, and controlling the execution respectively, we may use the following approximation to find the failure probability contribution from this minimal cut set:

$$Q_j \approx [q_P q_{CP} + \beta_{CP|P} \min(q_P, q_{CP})] [q_E q_{CE} + \beta_{CE|E} \min(q_E, q_{CE})] \quad (18)$$

Where we according to Table 1 assume that there are no common cause effects between planning and execution, and the β -values are found by equation (17).

8 Importance measures

There are a number of importance measures in the literature, where the Birnbaum's measure of reliability importance (see e.g., Rausand & Høyland, 2004) is one of the most commonly used. For a fault tree model Birnbaum's measure, $I^B(i)$, is defined as the change in the TOP-event probability, Q_0 as a function of the change in unreliability, q_i for component i , i.e., $I^B(i) = \partial Q_0 / \partial q_i$. In an event tree this is more complicated since there might be more than one "critical" end consequence. However, in the Risk_OMT project each of the event trees only has one critical end consequence, i.e., the leakage scenario. Thus the ordinary definition for $I^B(i)$ applies. The next challenge in the Risk_OMT project is that we have developed the model below the basic event level, i.e., the RIF structure. We then need to develop a Birnbaum like measure on the RIF-level. If we have a deterministic relation between a RIF, say RIF number j , and q_i we may apply the chain rule to find the importance measure. However, in the Risk_OMT model the RIFs are considered as random variables where observations (scores) and structural dependencies between the RIFs are used to express the uncertainty distribution regarding the RIFs. Since the RIFs are random variables and not parameters as in an ordinary fault tree or event tree, we need another definition of a "small change" in the value of the RIF. We propose to define a change in a RIF in terms of a shift in the *expected value* of the RIF. Let π_j be the posterior distribution of RIF j , and let ΔE_j be a (small) change in the expected value of π_j . Further assume that it is straight forward to establish a modified posterior, say π_j^Δ , given the shift in the expectation. Further, let F be the frequency of the critical end consequence, i.e., a leakage, where F depends on the posterior distribution of the RIFs, and in particular RIF j . A Birnbaum like measure for the importance of RIF j is then given by:

$$I_{\text{RIF}}^B(j) = [F(\pi_j^\Delta) - F(\pi_j)] / \Delta E_j \quad (19)$$

In order to implement the measure in equation (19) several aspects need to be considered. Shifting the posterior distribution of a RIF is not straightforward. The simplest situation is when we are considering first level RIFs, i.e., those influencing the parameters in the fault trees directly. It is rather easy to find the posterior distribution, π_j , from the BBN structure over the RIFs. If we next, approximate this distribution with for example a beta distribution, with some parameters we may rather easily find a new beta distribution (representing π_j^Δ) such that the expected value has changed by ΔE_j , and where we maintain the same variance in the distribution. In order to find $F(\pi_j^\Delta)$ we now just "disconnect" any parents of RIF j in the diagram to take complete control of the distribution of RIF j and then proceed with the quantification in a straight forward manner. Note that we need to make π_j^Δ discrete, which represents another approximation critical for the model in equation (19). In order to reduce the impact of this approximation it is therefore recommended to "disconnect" any parents of RIF j both when calculating $F(\pi_j^\Delta)$ and $F(\pi_j)$. The next challenge to treat is when we are working with the second level RIFs, i.e., those that only indirectly influences the fault tree parameters through the first order RIFs. We may still establish a shifted posterior distribution, π_j^Δ , but this will not have any impact on the scores for the first level RIFs. In the current implementation of the Risk_OMT model there is one, and only one second level RIF that influence each of the first level RIFs. Hence, it seems reasonable to shift the score of all the first level RIFs influenced by RIF j by a value ΔE_j . In this manner we get rid of the "momentum of the historical scores" in the model which would have damped the change in F as a function of ΔE_j for the second level RIFs.

9 Parameter estimation in the RISK_OMT model

The hybrid RISK_OMT model has several parameters that we need to specify:

q_L	Lowest value for the basic event probability, i.e., when all RIFs have state A
q_H	Highest value for the basic event probability, i.e., when all RIFs have state F
w_j	Standardized weight of level 1 RIF number j influencing the basic event
$V_{p,j}$	Structural importance of parent of level 1 RIF number j
$V_{s,j}$	Variance of the score of RIF number j given the true underlying RIF value r

In the analysis we distinguish between observations on the RIFs and observations on success or failure on the basic event level. The observations on the RIFs are typically scores assessed by interviews, surveys etc. These observations typically represent one installation for a period of time. Note that following one installation over time, the underlying true values of the RIFs are expected to change due to improvement projects etc. For a given period of time where the underlying RIFs are assumed to be more or less constant, we may have several observations related to the success or failure of the basic event considered. In principle, each time a maintenance activity is executed, we will have a new observation with respect to success or failure of the basic event.

9.1 Estimating $V_{s,j}$ and $V_{p,j}$

In order to estimate $V_{s,j}$ and $V_{p,j}$ we will utilize the values of the scores both for level 1 RIFs, and level 2 RIFs. We now recall the main assumptions and relations specified above:

- There is one and only one parent RIF (level 2 RIF) that points towards each child RIF (level 1 RIF)
- The score, S_j , of RIF number j is beta distributed with expected value r and variance $V_{s,j}$ given that the true underlying RIF value is r .
- Given the true underlying value p of the parent of level 1 RIF number j , this level 1 RIF is beta distributed with expectation p and variance $V_{p,j}$.

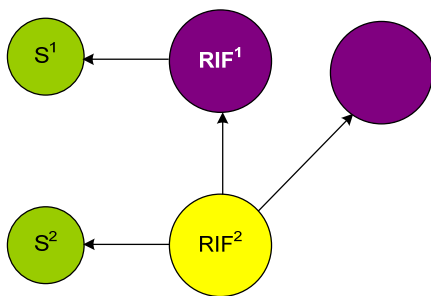


Figure 2 RIF structure for estimation

Figure 2 illustrates the principal situation for the estimation. To simplify the analysis, a first approach for the estimation would be to assume that the score of level 2 RIF represents the true underlying RIF. Now, let $s_{p,i}$ denote the value of the parent score of observation i . An observation here means that we have got a set of scores for one installation for a period of time where the RIFs are assumed to be relatively stable. Since we believe that the level 2 RIF is known, we may estimate parameters related to each level 1 RIF

independent of the other level 1 RIFs. Let $S_{c,i}$ denote the score of the level 1 RIF number j we are considering for observation i . Given the true value of the corresponding underlying level 1 RIF $R_{c,i} = r_{c,i}$, $S_{c,i}$ is beta distributed with expectation $r_{c,i}$ and variance $V_{S,j}$. Next, $R_{c,i}$ is beta distributed with expected value $s_{p,i}$ and variance $V_{p,j}$. We now apply the double expectation rule for the variance:

$$\text{Var}(X) = E(\text{Var}(X|Y)) + \text{Var}(E(X|Y)) \quad (20)$$

where X corresponds to the score of the child, $S_{c,i}$, and Y corresponds to the underlying level 1 RIF, $R_{c,i}$. Since $\text{Var}(X|Y) = \text{Var}(S_{c,i}|R_{c,i}) = V_{S,j}$, we have $E(\text{Var}(X|Y)) = V_{S,j}$. Further $E(X|Y) = E(S_{c,i}|R_{c,i}) = R_{c,i}$, hence $\text{Var}(E(X|Y)) = \text{Var}(R_{c,i}) = V_{p,j}$. Note that $\text{Var}(R_{c,i}) = V_{p,j}$ only if we know the value of the parent node, say $r_{p,i}$. Thus we have:

$$\text{Var}(S_{c,i} | r_{p,i} = s_{p,i}) = V_{S,j} + V_{p,j} \quad (21)$$

Equation (21) may now be used to estimate $\sigma^2 = V_{S,j} + V_{p,j}$. A natural estimator for σ^2 is:

$$\hat{\sigma}^2 = \frac{1}{n-1} \sum_{i=1}^n (S_{c,i} - s_{p,i})^2 \quad (22)$$

Note that we are not able to separate $V_{S,j}$ and $V_{p,j}$ from each other by using the empirical data. Hence, we need to utilize some expert judgements information. Note the different nature of the two variance terms. $V_{S,j}$ is a measure of how accurate the score is reflect the true underlying RIF. $V_{p,j}$ is a structural measure indicating how strong the influence the influence from the parent node on the child node is. It might be reasonable to assume that the variability of the score is less than the variability of the level 1 RIF as such, hence. i.e., $V_{S,j} < V_{p,j}$. Pragmatically we set if no other information is available $V_{S,j} = \frac{1}{4} V_{p,j}$. Also note that we in the estimation have assumed that the score of the parent RIFs equals the true underlying RIFs. We are also not able to extract the effect of this assumption in the model. In fact, we need to distribute $\hat{\sigma}^2$ not only to $V_{S,j} V_{p,j}$, but also to $V_{S,p}$, where index p here points to the parent node. Here, several child nodes might have the same parent node, hence a consistence check is required.

9.2 Estimating q_L , q_H and w_j

In order to estimate q_L , q_H and w_j we will only use the scores on the first level RIFs. A better approach would have been to insert the expected values of the underlying RIFs based on a combination of the score and the parent RIF, but for simplicity we just stick to the scores. The functional relation between the RIFs and the basic event probability is given by:

$$q = q_L \left(\frac{q_H}{q_L} \right)^{\sum_j w_j s_j} \quad (23)$$

We will first present a simple approach to estimate q_L , q_H and w_j by a transformation that brings equation (23) to a linear model by taking logarithms:

$$\ln q = \ln q_L + \sum_j w_j s_j \ln \left(\frac{q_H}{q_L} \right) = \beta_0 + \sum_j \beta_j s_j \quad (24)$$

where $\beta_0 = \ln q_L$ and $\beta_j = w_j \ln (q_H/q_L)$. For one combination of the score vector $\mathbf{s} = [s_1, s_2, \dots, s_r]$ there may be one or more observations. Typically there will be several observations if we have data on the basic event over a period of time where the score vector is assumed to remain constant, typically between surveys executed to assess the scores. We use i as an index to run through the relevant combinations of the score vector, and s_{ij} is the corresponding value of the score for level 1 RIF number j . Let x_i be the number of failures and n_i the number of “trials”, i.e., execution of the “basic” event. In the regression model of

equation (24) we would replace q with its estimate x_i/n_i . However, taking the logarithm on the left hand side will cause problems. This might be overcome with an empirical Bayesian approach where the prior distribution over q is set as the posterior distribution assuming all scores were the same, i.e., a beta distribution with $\alpha_0 = 1/2 + \sum_i x_i$ and $\gamma_0 = 1/2 + \sum_i n_i - \sum_i x_i$ (Jeffreys non informative prior). For observation number i the posterior distribution over q is then given by a beta distribution with parameters $\alpha_i = \alpha_0 + x_i$ and $\gamma_i = \gamma_0 + n_i - x_i$. The posterior mean is given by $\alpha_i/(\alpha_i + \gamma_i)$ which may be inserted in the logarithm in the left hand side of equation (24). We are now able to perform a standard linear regression. In the post processing we use that $\beta_0 = \ln q_L$ to find q_L . From the relation $\beta_j = w_j \ln (q_H/q_L)$ we easily find standardized weights, $w_j = \beta_j / \sum_j \beta_j$ and also $\ln (q_H/q_L) = \sum_j \beta_j$ such that an estimate for q_H might be obtained. The procedure is then as follows:

1. Find a prior distribution for q by either a direct approach, or use an empirical Bayesian approach where the parameters in the prior distribution are given by $\alpha_0 = 1/2 + \sum_i x_i$ and $\gamma_0 = 1/2 + \sum_i n_i - \sum_i x_i$.
2. For each observation find the Bayes estimate for q_i by calculating $\alpha_i = \alpha_0 + x_i$ and $\gamma_i = \gamma_0 + n_i - x_i$, and then calculate $\alpha_i/(\alpha_i + \gamma_i)$.
3. Calculate the left hand side of equation (24) by $Y_i = \ln q_i = \ln \alpha_i/(\alpha_i + \gamma_i)$.
4. Let x_{ij} be equal to the j^{th} level 1 RIF for observation i .
5. Equation (24) now reads $Y_i = \beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2} + \dots + \beta_r x_{ir} + \varepsilon_i$, where ε_i is the error term in the regression model.
6. Find the estimates by standard LS. Denote the estimates by β_j^* .
7. Find corresponding standardized weights $w_j^* = \beta_j^* / \sum_j \beta_j^*$, $j > 0$, $q_L^* = \exp(\beta_0^*)$, and $q_H^* = q_L^* \exp(\sum_j \beta_j^*)$

The procedure above may be conducted by simple calculation and a standard program for multiple linear regression (for example MS Excel or any statistical package). A standard maximum likelihood approach will require use of numerical methods for maximization of the likelihood function. The likelihood function is found by inserting q from equation (23) into the binomial probability distribution function for each observation yielding

$$L(q_L, q_H, w_1, \dots, w_r) = \prod_i \binom{n_i}{x_i} \left[q_L \left(\frac{q_H}{q_L} \right)^{\sum_j w_j s_{ij}} \right]^{x_i} \left[1 - q_L \left(\frac{q_H}{q_L} \right)^{\sum_j w_j s_{ij}} \right]^{n_i - x_i} \quad (25)$$

The maximization should be carried out under the constraint that the weights sum up to one.

References

- Aven, T., Sklet, S., and Vinnem, J.E. (2006) Barrier and operational risk analysis of hydrocarbon releases (BORA-Release); Part I Method description, *Journal of Hazardous Materials*, A137, 681-691
- Jeffreys, H. (1946). An Invariant Form for the Prior Probability in Estimation Problems. *Proceedings of the Royal Society of London. Series A, Mathematical and Physical Sciences* **186** (1007): 453–461 (http://en.wikipedia.org/wiki/Jeffreys_prior).
- Podofillini, L., V.N. Dang, P. Baraldi, M. Conti and E. Zio. 2009. A review of decision tree models for assessing Human Reliability Analysis dependence. Presented at ESREL2009, Prague, 7-10 September, 2009

- B.A. Gran, R. Bye, O.M. Nyheim, E.H. Okstad, J. Seljelid, S. Sklet, J. Vatn, J.E. Vinnem. 2012. Evaluation of the Risk OMT model for maintenance work on major offshore process equipment. *Journal of Loss Prevention in the Process Industries*, Volume 25, Issue 3, May 2012, Pages 582-593
- J.E. Vinnem, R. Bye, B.A. Gran, T. Kongsvik, O.M. Nyheim, E.H. Okstad, J. Seljelid, J. Vatn. 2012. Risk modelling of maintenance work on major process equipment on offshore petroleum installations. *Journal of Loss Prevention in the Process Industries*, Volume 25, Issue 2, March 2012, Pages 274-292